

ARTICLE

ChatGPT: A Case Study on Copyright Challenges for Generative Artificial Intelligence Systems

Nicola Lucchi 

Department of Law, University Pompeu Fabra, Barcelona, Spain.
Email: Nicola.Lucchi@upf.edu

Abstract

This article focuses on copyright issues pertaining to generative artificial intelligence (AI) systems, with particular emphasis on the ChatGPT case study as a primary exemplar. In order to generate high-quality outcomes, generative AI systems require substantial quantities of training data, which may frequently comprise copyright-protected information. This prompts inquiries into the legal principles of fair use, the creation of derivative works and the lawfulness of data gathering and utilisation. The utilisation of input data for the purpose of training and enhancing AI models presents significant concerns regarding potential violations of copyright. This paper offers suggestions for safeguarding the interests of copyright holders and competitors, while simultaneously addressing legal challenges and expediting the advancement of AI technologies. This study analyses the ChatGPT platform as a case example to explore the necessary modifications that copyright regulations must undergo to adequately tackle the intricacies of authorship and ownership in the realm of AI-generated creative content.

Keywords: Artificial intelligence; ChatGPT; copyright; data sharing; intellectual property; language models; training data

1. Introduction

News articles, academic papers, social media posts, photos and even chatbot chats are just some of the examples of how artificial intelligence (AI) is being put to use in the content creation process. Concerns have been voiced regarding the potential for AI to replace or mimic human behaviour as the technology continues to improve and find diverse applications across a wide range of sectors and industries. As a result, many organisations and academics in the field of law are starting to think about how AI might affect our society and the law.¹ Many areas of law are now grappling with the implications of these

¹ The emergence of generative AI infrastructures has presented new regulatory challenges in the field. For example, the European Commission is currently in the process of drafting the AI Act, the first law on AI by a major regulator, to regulate the emerging technology that has seen a surge in investments and popularity, particularly following the release of ChatGPT and its derivatives. Similar to the EU's General Data Protection Regulation (GDPR) in 2018, the EU AI Act has the potential to become a global standard, shaping the extent to which AI can have either positive or negative effects on individuals' lives worldwide. The draft is currently undergoing the trilogue phase, where EU parliamentarians and Member States will define the final details of the regulation. For more information, see the "Proposal for a Regulation of the European Parliament and of the Council Laying down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts", COM(2021) 206 final. See N Helberger and N Diakopoulos, "ChatGPT and the AI Act" (2023) 12 Internet Policy Review 10.14763/2023.1.1682.

technologies.² In this text, however, we will focus specifically on how AI-generated works may impact intellectual property law, with a particular emphasis on copyright law. In this work, we will briefly investigate some of the copyright issues linked with the usage of AI systems that recognise and generate text, known as large language models (LLMs),³ focusing specifically on the ChatGPT case study.⁴ Being a frequently utilised and well-known example of AI content production, ChatGPT provides a good lens for examining some of the fundamental copyright concerns at play in this rapidly growing sector.

ChatGPT is a language model created by OpenAI⁵ – a San Francisco-based AI company – that can generate replies in natural language to a variety of queries.⁶ A LLM is a highly effective type of machine learning process designed specifically for natural language processing tasks.⁷ Its main focus is on language modelling, which involves creating probabilistic models that can accurately predict the next word in a given sequence based on the preceding words.⁸ This is accomplished by training the model on large amounts of text data, which allows it to learn the probability of word occurrences and the patterns in language usage.⁹ The goal of language modelling is to create a system that can accurately generate human-like responses and recognise natural language input, making it an essential component of modern natural language processing applications.

It is important to stress that the language modelling task relies solely on form as training data, and therefore cannot inherently lead to the learning of meaning.¹⁰ These models are therefore characterized by their ability to “*agere sine intelligere*”¹¹; that is, to act without understanding exactly what they return as a result. This concept highlights the fascinating nature of their *modus operandi*, as they are able to perform complex tasks and produce results that can be remarkably accurate despite lacking a comprehensive understanding of the underlying processes. This phenomenon challenges conventional notions of intelligence, as these models have the potential to produce impressive results through a combination of sophisticated algorithms, vast amounts of data and intricate pattern recognition capabilities. Their ability to “*agere sine intelligere*” demonstrates the power of machine learning and its potential to revolutionise various fields, from natural language processing to image recognition and beyond. The advent of language models and various AI systems that produce content has been nothing short of a game-changer in

² An illustration of this is the recent decision by the Italian Data Protection Authority to take action against OpenAI's operations of ChatGPT in Italy, highlighting the tensions that exist between the EU's GDPR and the use of generative AI infrastructures that are trained on massive datasets containing both personal and non-personal data. See Garante per la Protezione dei Dati Personali, “Intelligenza artificiale: il Garante blocca ChatGPT. Raccolta illecita di dati personali. Assenza di sistemi per la verifica dell'età dei minori” (31 March 2023), <<https://www.garanteprivacy.it/home/docweb/-/docweb-display/docweb/9870847>> (last accessed 1 August 2023).

³ See, eg, Y Goldberg, *Neural Network Methods for Natural Language Processing* (Cham, Springer 2017) p 105; P Henderson et al, “Ethical Challenges in Data-Driven Dialogue Systems” (2018) *Proceedings of the 2018 AAAI/ACM Conference on AI, Ethics, and Society* 123; CD Manning et al, *An Introduction to Information Retrieval* (Cambridge, Cambridge University Press 2008) p 238.

⁴ The OpenAI GPT model was proposed in A Radford et al, “Improving Language Understanding by Generative Pre-Training” (2018) <<https://www.cs.ubc.ca/~amuham01/LING530/papers/radford2018improving.pdf>> (last accessed 1 August 2023).

⁵ See OpenAI <<https://openai.com/>>.

⁶ See OpenAI, “Introducing ChatGPT” <<https://openai.com/blog/chatgpt>> (last accessed 1 August 2023).

⁷ See EM Bender and A Koller, “Climbing towards NLU: On Meaning, Form, and Understanding in the Age of Data” in *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics* (Association for Computational Linguistics, Online, 2020) pp 5185–98 (defining the term “language model” as any system trained only on the task of string prediction, whether it operates over characters, words or sentences and sequentially or not).

⁸ Goldberg, *supra*, note 3, 105.

⁹ Manning et al, *supra*, note 3, 238.

¹⁰ Bender and Koller, *supra*, note 7, 5185.

¹¹ L Floridi, “AI as Agency without Intelligence: on Chat GPT, Large Language Models and Other Generative models” (2023) 36 *Philosophy & Technology* 1, 6.

today's world. These systems have the ability to generate text in any language, in any format and on any topic within seconds. The impact of these systems is therefore truly enormous, and it has given rise to numerous legal and ethical issues that need to be explored, especially from a copyright perspective.

Much of the current legal debate surrounding generative AI and copyright has focused on the potential protection of a “creative” product produced by AI technologies under copyright or similar intellectual property (here referred to as “the output”).¹² However, it is important to recognise that there are also significant copyright issues associated with the use of copyrighted information to train and develop AI systems (here referred to as “the input”). Indeed, AI systems require massive amounts of training data, which frequently contain copyrighted information, in order to create high-quality outputs. This raises concerns about whether and how such data may be collected and utilised lawfully, as well as concerns about derivative works¹³ and fair use.¹⁴ Furthermore, as AI systems grow more prevalent and vital in our daily lives, it is critical to address the copyright challenges arising from the process of training AI models. This involves the creation of derivative works from protected sources, often requiring modifications or manipulations of data to enhance their suitability for training purposes. Recently, the legal debate surrounding AI has intensified, leading to numerous lawsuits against creators of generative AI systems such as ChatGPT, alleging copyright infringement.¹⁵ These lawsuits raise legitimate concerns about the unauthorised use of copyrighted material in order to create new creative content. In light of these challenges, a comprehensive and holistic approach is needed to tackle the copyright problems associated with AI, considering both the inputs and outputs of AI systems. This investigation will also delve deeper into the policy rationales for considering a free or open-access approach to AI training data, with the goal

¹² Many scholars have given attention to the question of ownership and authorship: see, eg, R Abbot, *The Reasonable Robot* (Cambridge, Cambridge University Press 2020); E Bonadio and L McDonagh, “Artificial Intelligence as Producer and Consumer of Copyright Works: Evaluating the Consequences of Algorithmic Creativity” (2020) 2 Intellectual Property Quarterly 112; E Bonadio et al, “Intellectual property aspects of robotics” (2018) 9 European Journal of Risk Regulation 655; A Bridy, “Coding Creativity: Copyright and the Artificially Intelligent Author” (2012) 5 Stanford Technology Law Review 1; R Denicola, “Ex Machina: Copyright Protection for Computer-Generated Works” (2016) 69 Rutgers University Law Review 251; T Dornis, “Artificial Creativity: Emergent Works and the Void in Current Copyright Doctrine” (2020) 22 Yale Journal of Law & Technology 1; J Grimmelmann, “There’s No Such Thing as a Computer-Authored Work – And It’s a Good Thing, Too” (2016) 39 Columbia Journal of Law & the Arts 403; A Guadamuz, “Do Androids Dream of Electric Copyright? Comparative Analysis of Originality in Artificial Intelligence Generated Works” (2017) 2 Intellectual Property Quarterly 169; AH Khoury, “Intellectual Property Rights for Hubots: On the Legal Implications of Human-like Robots as Innovators and Creators” (2017) 35 Cardozo Arts & Entertainment Law Journal 635; M Lemley and B Casey, “Remedies for Robots” (2019) 86 University of Chicago Law Review 1311; E Bonadio and N Lucchi (eds), *Non-Conventional Copyright: Do New and Non Traditional Works Deserve Protection?* (Cheltenham, Edward Elgar 2018); D Lim, “AI & IP: Innovation & Creativity in an Age of Accelerated Change” (2018) 52 Akron Law Review 813; P Samuelson, “Allocating Ownership Rights in Computer-Generated Works” (1986) 47 University of Pittsburgh Law Review 1185; P Yu, “The Algorithmic Divide and Equality in the Age of Artificial Intelligence” (2020) 72 Florida Law Review 331; R Yu, “The Machine Author: What Level of Copyright Protection Is Appropriate for Fully Independent Computer-Generated Works?” (2017) 165 University of Pennsylvania Law Review 1245; DL Burk, “Thirty-Six Views of Copyright Authorship, by Jackson Pollock” (2020) 58 Houston Law Review 263.

¹³ In copyright law, a derivative work is a work that is based on one or more pre-existing works, such as a translation, adaptation, sequel or a work that is based on another work in some way. A derivative work is considered to be a new work, but it still retains some of the characteristics of the original work. A work can serve a transformative purpose, even if it does not alter the content of the original work. See RA Reese, “Transformativeness and the Derivative Work Right” (2008) 31 Columbia Journal of Law & the Arts 467, 485.

¹⁴ 17 U.S.C. § 107 (2018).

¹⁵ See, eg, *Getty Images (US), Inc. v. Stability AI, Inc.*, No. 1:23-cv-00135-GBW (D. Del. Mar. 29, 2023); *Silverman et al. v. OpenAI, Inc. et al.*, No. 4:23-cv-03416 (N.D. Cal. Jul. 7, 2023); *Tremblay et al. v. OpenAI, Inc. et al.*, No. 4:2023-cv-03223 (N.D. Cal. Jul. 7, 2023). Essentially, all of these very recent lawsuits (still pending) allege that the incorporation of training data by generative AI models is an infringement of copyright holders’ rights.

of potentially proposing legislation that encourages the responsible and ethical use of such data while protecting intellectual property rights.¹⁶

In light of the scenarios mentioned, this article aims to propose effective strategies that can address the legal issues arising from AI system development while simultaneously safeguarding the rights of copyright holders and competitors. Given the rapid advancements in AI technology, it is essential to establish a robust legal framework as well as a set of rules to ensure the protection of all stakeholders involved.

The article is divided into three parts. In the first part, we set the stage by discussing the capabilities and limitations of powerful language models, including their potential and actual applications as well as their limitations. In the second part, we look at the case study of ChatGPT and explore how this generative AI system works, discussing specific copyright concerns. In particular, we explore the practical applications of ChatGPT-generated text and address important issues related to ownership and copyright, especially when the content is created by a machine rather than a human author. In the third part, we further analyse the ChatGPT case study by focusing on the challenges related to training data and copyright. We address the complexities of data ownership and use and explore the different types of data used to train the ChatGPT models. We also analyse some recent court cases and examine the ethical and legal dilemmas that arise when dealing with large datasets. In addition, we investigate and evaluate a number of potential alternatives that can effectively safeguard copyrighted training data used in the field of AI to feed generative AI systems. To conclude, we provide a concise summary of our comprehensive analysis and highlight the significant findings and insights we have gained from investigating ChatGPT language models. We acknowledge the obstacles and limitations that must be overcome to advance these models and emphasise the importance of addressing these issues responsibly and ethically.

Due to the breadth of the legal and policy issues presented in this article, and in order to support the “circulation of models”, the article will provide a mix of civil and common law elements based on the regulations and laws of multiple jurisdictions, with a particular emphasis on the USA, the UK and the European Union (EU).

II. ChatGPT: the AI language model revolutionising content creation

You may have heard about ChatGPT, which attracted millions of users quickly after its launch. OpenAI released ChatGPT as a chatbot¹⁷ designed primarily to assist with online customer care. Yet, the AI technology underpinning this program has outperformed all forecasts and grown in popularity.

¹⁶ Numerous policy initiatives have been implemented at both the international and EU levels pertaining to the intersection of AI and copyright (even if they mainly focused on the question of the output of AI). See, eg, European Parliament, Resolution on a comprehensive European industrial policy on artificial intelligence and robotics, (2018/2088 (INI)), 12 February 2019; European Parliament resolution of 20 October 2020 on intellectual property rights for the development of artificial intelligence technologies (2020/2015 (INI)); WIPO, “WIPO Technology Trends 2019: Artificial Intelligence” <<https://www.wipo.int/publications/en/details.jsp?id=4386>> (last accessed 1 August 2023); WIPO, “Revised Issues Paper on Intellectual Property Policy and Artificial Intelligence”, 2020 (WIPO/IP/AI/2/GE/20/1 REV) <https://www.wipo.int/edocs/mdocs/mdocs/en/wipo_ip_ai_2_ge_20/wipo_ip_ai_2_ge_20_1_rev.pdf> (last accessed 1 August 2023).

¹⁷ Chatbots are computer programs that use natural language processing (NLP) and machine learning algorithms to understand and interpret user input and respond with appropriate pre-programmed messages or actions designed to simulate human-to-human conversation, typically via text-based interfaces such as messaging apps, websites or mobile applications. See R Dale, “The return of the chatbots” (2016) 22 *Natural Language Engineering* 811; E Adamopoulou and L Moussiades, “An Overview of Chatbot Technology” (2020) 584 *Artificial Intelligence Applications & Innovations* 373.

What is the definition of ChatGPT? We have chosen to ask the question directly to the chatbot:

ChatGPT is an AI language model developed by OpenAI, which uses Natural Language Processing (NLP) to generate human-like text in response to various inputs. It is essentially an auto-generating chatbot that extracts data from various sources, processes it, and produces grammatically correct and contextually appropriate responses. ChatGPT can be used for a variety of applications, including customer service, chatbots, and virtual assistants. It has been trained on a massive amount of text data and is constantly learning and improving.

The ChatGPT software application was designed to facilitate and support interactive communication by generating automatic text-based responses. The methodology involves integrating information from various sources, such as books, journals, websites and articles, resulting in the development of original and engaging discourse. The field of natural language processing (NLP) involves the use of computational methods to facilitate the understanding and analysis of large amounts of data in a way that mimics human language. Computer systems are capable of acquiring significant amounts of data or executing commands by utilising NLP functionalities.¹⁸ It is worth noting that the operational functionalities of ChatGPT are limited to the understanding of syntax. The employed methodology involves a comprehensive analysis of the sentence's syntactic structure, along with the identification of the interdependencies among its constituent elements. Although ChatGPT has undergone comprehensive training in handling vast datasets, it continues to depend on statistical techniques for analysis and comprehension.¹⁹ The aforementioned statement implies that the system possesses the ability to identify patterns in the utilisation of specific terminologies and idiomatic expressions and employs these patterns to generate informed assumptions regarding the intended interpretation of a provided sentence or phrase. It is crucial to recognise that this statement precludes any inferences about ChatGPT's understanding of the aforementioned principles. In the event that an individual were to inquire of ChatGPT regarding the essence of existence, the response provided may appear to be lucid and perceptive; however, it may lack a comprehensive comprehension of the philosophical principles and hypotheses that underlie the inquiry. ChatGPT is just an advanced chatbot that employs NLP to comprehend vast quantities of information and produce responses that closely resemble human language. Its capabilities are currently limited to syntactic understanding, meaning that it can analyse the grammatical structure of sentences and comprehend how words and phrases relate to each other within a sentence. Although ChatGPT is able to detect patterns in the use of words and phrases in context and to use this information to make educated guesses about the meaning of a sentence or phrase, it still approaches semantic understanding through statistical analysis.²⁰ For this reason, ChatGPT's responses

¹⁸ See generally D Jurafsky and JH Martin, *Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition* (Upper Saddle River, NJ, Pearson 2009). See also A Guadamuz, "Authors sue OpenAI for copyright infringement" (*TechnoLlama*, 8 July 2023) <<https://www.technollama.co.uk/authors-sue-openai-for-copyright-infringement>> (last accessed 1 August 2023; explaining how large language model are trained).

¹⁹ See TB Brown et al, "Language Models Are Few-Shot Learners" (2020) arXiv preprint <<http://arxiv.org/abs/2005.14165>> (last accessed 1 August 2023; noting that despite being trained on massive amounts of data, these models still lack true semantic understanding and instead approximate it through statistics).

²⁰ "Syntactic understanding" refers to the ability to understand the structure and rules of language, including grammar and syntax. "Semantic understanding", on the other hand, refers to the ability to understand the meaning of language, including concepts and context. See CD Manning and H Schutze, *Foundations of Statistical Natural Language Processing* (Cambridge, MA, MIT Press 1999) p 3.

do not always reflect true understanding of the underlying concepts or theories related to a particular question. ChatGPT is capable of independently stimulating dialogues and thus has the potential to produce content protected by intellectual property, including but not limited to articles, music lyrics, programming codes and text translations. The results produced by ChatGPT depend on the data it was programmed with and the computational techniques used and may not always be suitable for all targets. While acknowledging the commendable nature of ChatGPT, it is important to point out that its AI-powered functions are created without human intervention. Even though AI has remarkable precision, it is not free of limitations. Therefore, it is important that individuals – in order to prevent potential problems or errors – review and modify the designs to ensure that they meet established standards for accuracy and efficiency for specific usage scenarios.

III. An analysis of the ChatGPT case study: the question of the output

The emergence and widespread application of AI systems in the creative sectors have raised concerns about the rightful ownership of intellectual property and the protection of copyright. To gain a more comprehensive understanding of the above issues, our research has specifically focused on ChatGPT, a language acquisition model developed by OpenAI. ChatGPT – as mentioned earlier – is a generative AI tool. A full understanding of the complexity of the evolution and ownership of the ideas created by AI can be achieved by evaluating the results and sharing observations about the interaction between humans and computer systems. This emphasises the importance of a recasting of legal rules in order to adequately deal with this set of challenges. In this perspective, the ChatGPT case study can be used to initiate a debate about the legal and ethical implications of rapid technological advancements, as well as the concerns associated with the application of AI in the creative industries. The use of generative techniques such as ChatGPT raises significant issues regarding intellectual property, authorship and the scope of copyright protection for material created with generative AI systems. These aspects have attracted considerable attention in the legal field.²¹ The primary inquiry relates to the need to determine the rightful owner of copyright in content generated by AI, whether a natural person or a legal entity. We attempted to ask ChatGPT directly about this question, and the platform provided us with the following response:

As an AI language model, I do not own the copyright of the text generated with my help. The ownership of the text belongs to the user who inputs the prompts and generates the output.²²

Of course, the question of who owns the content developed by ChatGPT is more complicated and may require further clarification or reference to be fully answered. Because ChatGPT is an AI system that generates results based on training data and user input, it is difficult to identify specific authors. Authorship of AI-generated content may depend on a variety of factors, including the purpose of the content, the intent of the user and the legal framework at the time.

In the case of ChatGPT, the author of the content may be based around the individual who created the prompt or input for the response. If the user provides the input, they can take ownership of the output. If the input comes from ChatGPT's training data or other

²¹ See *supra*, note 12.

²² K Anderson, "ChatGPT says it's not an author" (*The Geyser*, 13 January 2023) <<https://www.the-geyser.com/chatgpt-says-its-not-an-author/?ref=the-geyser-newsletter>> (last accessed 1 August 2023).

sources, it may be more difficult to identify the owner. Usually, the copyright owner of texts created with tools such as ChatGPT is the person or organisation that provided the original ideas and data on which the system is based or the person who creatively implemented the instructions in the prompt.

As an AI language model, the text generated is not, per se, protected by copyright law, as copyright law generally recognises the human creator of an original work as the copyright owner. However, in some cases, the text generated may be considered original enough to be protected by copyright if it was created with sufficient human input or intervention. For example, the resulting work could be deemed sufficiently unique to be protected by copyright law if someone uses the replies as a starting point and then adds significant creative or original content, such as editing, adding commentary or analysis or merging it into a bigger work. The individual who added the extra creative or original content in this scenario would normally own the copyright for the final product.

AI's inability to hold copyright stems from its legal identity or status as a non-human entity.²³ While the Berne Convention and other international copyright regulations do not require human authorship, many countries, such as the USA and those in the EU, place importance on the presence of a human being as the creator of a work.²⁴ In addition, copyright law itself adopts a predominantly anthropocentric approach, as exemplified by the copyright term "70 years after the calendar year in which the author of the work died". This term inherently assumes that the author is a human being, subject to mortality.

Original pieces of art produced with AI assistance or by automated means are not novel occurrences. Some might argue that what we are witnessing with AI systems is simply the repetition of history. After all, copyright laws have always had to evolve and keep pace with emerging technologies and their effects on society. This pattern can be observed, for instance, with the arrival of photography, motion pictures, computer programs and various other novel forms of creative expression.²⁵ A historical example of this is the well-known case in *Burrow-Giles v. Sarony* in 1885,²⁶ which was argued before the US Supreme Court. At issue was whether a photograph could be considered a copyrighted work, given that the image was created by a camera rather than a human being. In its ruling, the Court held that the photographer, who was the person behind the camera, was the author of the photograph and therefore had exclusive copyright over it. This rationale persisted even in

²³ A Guadamuz, "Artificial Intelligence and Copyright" (*World Intellectual Property Office Magazine*, October 2017) <https://www.wipo.int/wipo_magazine/en/2017/05/article_0003.html> (last accessed 1 August 2023; listing many jurisdictions requiring human authorship for copyright protection).

²⁴ See, eg, Case C-5/08 *Infopaq International A/S v Danske Dagblades Forening*; Case C-145/10 *Eva-Maria Painer v Standard Verlags GmbH and Others* (emphasising that a work achieves the originality criterion when it contains a "intellectual production" with the author's own creative visible touch). Similarly, in the US – shortly after *Naruto v. Slater*, No. 16-15469 (9th Cir. 2018) – the *Compendium of US Copyright Practices* was updated to provide that "the US Copyright Office will register an original work of authorship, provided that the work was created by a human being" and "will refuse to register a claim if it determines that a human being did not create the work." See US Copyright Office, *Compendium of US Copyright Office Practices* section 306, 3rd edition (as emended 2021). See also the recent rejection of the US Copyright Office to register the copyright of a picture generated by an AI: "Second Request for Reconsideration for Refusal to Register a Recent Entrance to Paradise" (Correspondence ID 1-3ZPC6C3; SR # 1-7100387071), U.S. Copyright Off. Rev. BD. 1-2 (14 February 2022) <<https://www.copyright.gov/rulings-filings/review-board/docs/a-recent-entrance-to-paradise.pdf>> (last accessed 1 August 2023). After receiving the Office's final rejection, legal action was taken, resulting in a lawsuit against the US Copyright Office. In a very recent development, the US District Court for the District of Columbia ruled on the case. The court's decision, as seen in *Thaler v. Perlmutter*, No. 22-1564 (D.D.C. Aug. 18, 2023), declared that AI-created works are ineligible for protection under US copyright law.

²⁵ See eg B Sherman and L Wiseman (eds), *Copyright and the Challenge of the New* (Alphen aan den Rijn, Kluwer Law International 2012) p 1 (noting that "one of the most challenging things about copyright law is that it is constantly subject to change").

²⁶ See *Burrow-Giles Lithographic Co. v. Sarony*, 111 U.S. 53 (1884).

cases where the machine was responsible for most of the work, since it was recognised that human input was necessary and indispensable for the creation of the work.²⁷ The UK is the only country where the concept of “computer-generated works” is recognised in domestic law, and it seeks to address the issue in a practical way by expanding the concept of authorship. Section 9(3) of the Copyright, Designs and Patents Act 1988 (CDPA) provides that the person who makes the necessary arrangements for the creation of the work shall be deemed to be the author.²⁸ However, determining who is the “arranger” is not always easy and often needs to be determined on a case-by-case basis. In particular, according to contemporary standards for defining “computer-generated” works, the UK approach can be considered quite outdated.²⁹ These provisions became law in 1988, and the AI systems available today are vastly different from the computer systems that existed at that time. Moreover, given the complexity of modern AI programming, there is considerable uncertainty in determining which party is responsible for the “arrangements necessary for the creation of the work”.³⁰

Being an AI language model, ChatGPT also lacks legal identity and the ability to possess property or assets in the conventional sense because it is not a human. Even if the content created by an AI language model is original and creative enough to be protected by copyright law, the AI will not own it. According to the various jurisdictions mentioned, the copyright for the material could belong to the individual or entity that has legal authority over the AI, such as the AI system’s developer or owner. In some instances, the content’s copyright may belong to the human users who contributed to or edited the AI-generated work.

The practical approach, then, is to grant copyright to the people behind the machines, namely the programmer, the user and the owner. These key actors are the human or human-owned entities behind the process of AI production and, accordingly, the actors at the centre of the legal discussion about copyright in AI-generated works. On the other hand, if we consider only the statutory and common law understanding of the doctrine of originality and the requirement of human authorship, there is indeed no copyright in works created by AI, and copyright-free works naturally belong to the public domain.

IV. An analysis of the ChatGPT case study: the question of originality

The issue of authenticity of the output of AI generative systems presents another challenge and variation, especially with respect to ChatGPT. While chatbots are excellent at generating responses that engage humans in conversation, these responses run the risk of being unoriginal, completely invented and simply repeating information from the past. The use of chatbots and generative tools for content creation can lead to problems, especially in cases where the resulting output requires distinctiveness and appeal.

²⁷ See R Yu, “The Machine Author: What Level of Copyright Protection Is Appropriate for Fully Independent Computer-Generated Works?” (2017) 165 *University of Pennsylvania Law Review* 1245, 1253.

²⁸ See § 9(3) of the UK Copyright, Designs and Patents Act 1988. For a more detailed discussion, see E Bonadio et al, “Will Technology-Aided Creativity Force Us to Rethink Copyright’s Fundamentals? Highlights from the Platform Economy and Artificial Intelligence” (2022) 53 *International Review of Intellectual Property and Competition Law* 1174, 1187.

²⁹ See Intellectual Property Office (UK), “Artificial intelligence call for views: copyright and related rights” (UK Government, 2020) <<https://www.gov.uk/government/consultations/artificial-intelligence-and-intellectual-property-call-for-views/artificial-intelligence-call-for-views-copyright-and-related-rights>> (last accessed 1 August 2023).

³⁰ For some additional critical comments on this provision, see PB Hugenholtz and JP Quintais, “Copyright and Artificial Creation: Does EU Copyright Law Protect AI-Assisted Output?” (2021) 52 *International Review of Intellectual Property and Competition Law* 1190, 1211 (noting that “since the introduction of the regime on computer-generated works in UK law in 1988, this has led to just a single court decision, which has not clarified this issue”).

Intellectual property ownership is closely related to the uniqueness of AI chatbots. As AI production processes become more sophisticated and generate content that closely resembles human-created content, it is important to establish clear guidelines and rules for creation and submission. The significance of this matter is particularly salient in domains such as journalism and creative writing, given that the utilisation of AI-generated content can engender ethical and legal dilemmas. Typically, the degree to which creative works are safeguarded is contingent upon their level of uniqueness. While the Berne Convention does not expressly mandate the condition of “originality” for copyrighted works, several nations enforce this prerequisite.³¹ As a consequence, the originality requirement is a prerequisite for the granting of copyright protection to literary, dramatic, musical and artistic works. Currently, the prevailing approach to determining whether a work is original is by evaluating whether it is “the author’s own intellectual creation”.³² This means that the work must have an intellectual content that goes beyond the mere combination of its individual parts, taking into account the overall impression. However, the standard for originality varies across different jurisdictions. For instance, US law adopts the “minimal degree of creativity” test, which was established in the *Feist v. Rural* case,³³ while the EU requires that the work be an author’s own “intellectual creation”.³⁴ But the issue of originality becomes more nuanced when considering content created by AI. Copyright laws may apply to certain AI-generated material if it is made with sufficient human input or participation. An AI chatbot may be regarded to have created an original work under copyright law if a human offers input or instruction to the bot to create a particular work, such as a tale or a song, and the bot then develops the final output based on that input. Some AI-generated content, however, may be less unique than expected and more derivative or based on previously published works. For example, if AI-generated systems merely reproduce data or existing information without adding significant ideas or original content, the output they create cannot be considered unique enough to be copyright protected. By applying the US “minimal degree of creativity” test for originality, which sets a low bar for copyright protection, one could argue that ChatGPT’s output meets this standard. This is because ChatGPT utilises sophisticated NLP techniques to generate text that is not merely a repetition of its input data, indicating some level of creativity. However, under the EU’s standard for originality, AI-generated works may not qualify, as they lack the creative choices and personal expression of a human author.

ChatGPT, as an AI language model, has undergone extensive training using vast amounts of textual data gathered from diverse online and offline sources. By leveraging these training data, it generates responses to user queries encompassing a broad spectrum of subjects, ranging from common knowledge to specialised fields. When presented with a

³¹ See E Bonadio and N Lucchi, “Introduction: Setting the Scene for Non-Conventional Copyright” in E Bonadio and N Lucchi (eds), *Non-Conventional Copyright: Do New and Atypical Works Deserve Protection?* (Cheltenham, Edward Elgar Publishing 2018) p. 6.

³² The Court of Justice of the European Union (CJEU) provided that in order to pass the threshold for copyright protection, a “work” needs to be “original” in the sense that it is the author’s “own intellectual creation”. See Case C-05/08, *Infopaq International v. Danske Dagblades Forening* (2009) ECLI:EU:C:2009:465 (Infopaq). This ruling has been reaffirmed in subsequent CJEU cases, including *Levola Hengelo*, *Funke Medien*, *Cofemel* and *Brompton Bicycle*. See the following case references for more information: Case C-310/17 (*Levola Hengelo*); Case C-469/17 (*Funke Medien NRW GmbH v. Bundesrepublik Deutschland* (2019) ECLI:EU:C:2019:623); Case C-683/17 (*Cofemel - Sociedade de Vestuário SA v. G-Star Raw CV* (2019) ECLI:EU:C:2019:721); and Case C-833/18 (*SI and Brompton Bicycle Ltd v. Chedech/Get2Get* (2020) ECLI:EU:C:2020:461).

³³ See *Feist Publications, Inc., v. Rural Telephone Service Co.*, 499 U.S. 340 (1991).

³⁴ See Case C-5/08, *Infopaq Int’l A/S v. Danske Dagblades Forening*, 2009 E.C.R. I-6569 (setting out the EU originality standard for copyright protection). According to the court, copyright can only protect the author’s own intellectual creativity, which reflects their personality and has a certain originality. It states that originality is characterised by the author’s individuality and expression and that the work must reflect the author’s personal touch and creativity.

question, ChatGPT examines the overall context and relevant keywords to formulate a response based on learned relationships and patterns derived from the training data. These responses are algorithmically generated and do not rely on the respondent's personal opinions or experiences. It is important to note again that ChatGPT establishes statistical correlations between words without genuine comprehension of the underlying meaning. The tool excels at producing high-quality written content across various domains, saving considerable time compared to human effort, thanks to its extensive database and syntactic correlation capabilities. However, due to the absence of human authorship, ChatGPT lacks the necessary human creative input required to substantiate a copyright claim.³⁵ In copyright law, in fact, the act of creating a copyrightable work is typically associated with human creativity and authorship. Therefore, if there is no human involvement in the creation process, there is a lack of originality, and consequently the work may not be eligible for copyright protection. The concept of personhood is crucial in this context, as it distinguishes between entities with naturalistic dimensions of life and self-awareness and those that do not possess these attributes. Robots or other AI technologies, regardless of their level of autonomy, cannot be considered as persons under ethical and legal frameworks. The qualification of personhood plays a significant role in copyright law as it serves as a boundary for attributing creative authorship and the associated rights. This distinction is based on the understanding that copyright protection is intended to incentivise and reward the unique and subjective contributions of human creators.

While robots and AI systems can generate content or imitate human-like behaviours, they lack the essential qualities that define personhood, such as consciousness, intentionality and the capacity for subjective experience. These intrinsic limitations prevent us from assimilating them into the category of persons within ethical and legal contexts.

Therefore, it remains firmly established that copyright protection requires a human element, where the creative efforts and expression originate from individuals possessing the characteristics and attributes inherent to personhood.

In light of this analysis, it can be deduced that in order for an AI system to fully replace a human author, it would require the capacity to independently conceptualise and complete a creative work without relying on explicit training or pre-programmed instructions.³⁶ As technology continues to progress, it is reasonable to envisage a gradual reduction in human involvement in the creative process, leading to the emergence of new artistic creations that cannot be attributed to a specific or recognised artist.³⁷ While the current capabilities of ChatGPT may not align with this vision, a completely revolutionary future appears to be within reach.

V. An analysis of the ChatGPT case study: the question of the input

At present, the intellectual property discourse pertaining to AI predominantly revolves around the issues regarding the authorship and creative ownership of the outcomes

³⁵ There is a common consensus on this: recently, the US Copyright Office has issued an official policy declaration that firmly denies the registration of copyright for works produced entirely by AI without any human involvement. Copyright Office, "Copyright Registration Guidance: Works Containing Material Generated by Artificial Intelligence" (16 March 2023) 37 CFR Part 202 <<https://www.govinfo.gov/content/pkg/FR-2023-03-16/pdf/2023-05321.pdf>> (last accessed 1 August 2023).

³⁶ See JC Ginsburg and LA Budiardjo, "Authors and Machines" (2019) 34 Berkeley Technology Law Journal 343.

³⁷ See, eg, RB Abbott and E Rothman, "Disrupting Creativity: Copyright Law in the Age of Generative Artificial Intelligence" Florida Law Review, forthcoming, available at SSRN <<https://ssrn.com/abstract=4185327>> (last accessed 1 august 2023; claiming that AI should be recognised as an author when it performs tasks equivalent to human authors).

generated by AI systems. Despite the ongoing discourse, there appears to be a significant gap in the examination of the legal questions that emerge in the management of intellectual property rights pertaining to the inputs, namely the data employed in the training of these AI systems.³⁸ The second and more fundamental question to be addressed here is whether the use of copyrighted material to train generative AI programs represents an infringement of copyright. Indeed, machine learning heavily relies on vast amounts of training data to achieve accurate results, including in facial recognition, stop sign recognition, natural language recognition and translation generation. This is especially important when it comes to ChatGPT because it relies on large amounts of training data being fed into the system.³⁹ In order to create interactive and authentic articles, ChatGPT needs to ingest information, including text, images and other content, from publicly available websites on the Internet.

To facilitate the training of AI algorithms, various techniques are used, including text and data mining (TDM)⁴⁰ as well as generative deep learning techniques.⁴¹ TDM processes involve the extraction and analysis of vast amounts of data to identify meaningful insights and patterns, which can then be leveraged to improve the performance of AI models.⁴² TDM has become an essential tool in the field of AI, enabling researchers and data scientists to explore vast amounts of unstructured data and extract valuable information that would otherwise be impossible to obtain manually. By analysing these vast amounts of data, AI algorithms can learn from these patterns and make predictions with a high degree of accuracy, facilitating the creation of content, discoveries and innovations. So, without access to large volumes of data, AI algorithms would struggle to “learn” and improve their performance. Therefore, it is clear that the future of AI hinges on TDM and its capacity to extract and analyse data on a large scale. However, a significant challenge lies in the fact that AI systems cannot learn from art in the same way humans do, since they require an exact copy of the artwork in their training dataset.⁴³ This necessitates the creation of a

³⁸ Scientific interest in this topic is growing among legal scholars: see, eg, A Strowel, “ChatGPT and Generative AI Tools: Theft of Intellectual Labor?” (2023) 54 *International Review of Intellectual Property and Competition Law* 491; G Franceschelli and M Musolesi, “Copyright in generative deep learning” (2022) 4 *Data & Policy* e17; E Bonadio et al., “Can Artificial Intelligence Infringe Copyright? Some Reflections”, in R Abbott (ed.), *Research Handbook on Intellectual Property and Artificial Intelligence* (Cheltenham, Edward Elgar 2022); J Quang, “Does Training AI Violate Copyright Law?” (2021) 36 *Berkeley Technology Law Journal* 1407; B Sobel, “A Taxonomy of Training Data: Disentangling the Mismatched Rights, Remedies, and Rationales for Restricting Machine Learning”, in R Hilty et al (eds), *Artificial Intelligence and Intellectual Property* (Oxford, Oxford University Press 2021) pp 221–42; G Abbamonte, “The rise of the artificial artist: AI creativity, copyright and database right” (2021) 43 *European Intellectual Property Review* 702; MA Lemley and B Casey, “Fair Learning” (2021) 99 *Texas Law Review* 743; P Keller, “Protecting creatives or impeding progress? Machine learning and the EU copyright framework” (*Kluwer Copyright Blog*) <<https://copyrightblog.kluweriplaw.com/2023/02/20/protecting-creatives-or-impeding-progress-machine-learning-and-the-eu-copyright-framework/>> (last accessed 1 August 2023); J Vincent, “The scary truth about AI copyright is nobody knows what will happen next” (*The Verge*, 15 November 2022) <<https://www.theverge.com/23444685/generative-ai-copyright-infringement-legal-fair-use-training-data>> (last accessed 1 August 2023); CJ Craig, “The AI-Copyright Challenge: Tech-Neutrality, Authorship, and the Public Interest” (14 December 2021) Osgoode Legal Studies Research Paper <<https://ssrn.com/abstract=4014811>> (last accessed 1 August 2023); R Abbott (ed.), *Research Handbook on Intellectual Property and Artificial Intelligence* (Cheltenham, Edward Elgar 2022).

³⁹ See Brown et al, *supra*, note 19; A Radford et al, “Language Models Are Unsupervised Multitask Learners” (2019) 8 *OpenAI Blog* 1 <<https://insightcivic.s3.us-east-1.amazonaws.com/language-models.pdf>> (last accessed 1 August 2023); J Devlin et al, “BERT: Pre-Training of Deep Bidirectional Transformers for Language Understanding” (2018) *ArXiv:1810.04805* [Cs], arXiv.org <<http://arxiv.org/abs/1810.04805>> (last accessed 1 August 2023).

⁴⁰ See E Alpaydın, *Introduction to Machine Learning* (Cambridge, MA, MIT Press 2004) p 2 (explaining that TDM is an essential tool for machine learning and data mining, particularly in cases where the data are too numerous or too complex for humans to analyse manually).

⁴¹ Franceschelli and Musolesi, *supra*, note 38.

⁴² Alpaydın, *supra*, note 40.

⁴³ Lemley and Casey, *supra*, note 38, 775.

training set of millions of examples by making copies of copyrighted images, videos, audio or text-based works. Consequently, the question of whether machine copying should fall under fair use or other copyright exceptions arises. On the other hand, we have generative deep learning, a specialised branch of deep learning that focuses primarily on the task of generating novel data.⁴⁴ Generative models are crucial in this domain as they provide a probabilistic framework for describing the data generation process.⁴⁵ By harnessing these models, it becomes possible to generate new data samples through the process of sampling. These techniques employ deep neural networks, which are artificial neural networks with multiple layers, to learn and replicate the patterns, structures and statistical properties present in the training data. ChatGPT is precisely a form of generative deep learning technique that harnesses the power of deep learning models, particularly the GPT (Generative Pre-trained Transformer) architecture.

The problem here is that established companies such as Google, Facebook, Amazon and OpenAI have access to large collections of language and image data, which they can use for AI purposes.⁴⁶ Access to large collections of language and image data can be also considered a competitive advantage in the field of AI.⁴⁷ As a consequence, these companies can leverage their existing datasets to train and develop more advanced AI models, which in turn can improve their products and services. This can create a legal problem for new entrants because the ownership and licensing of datasets can be complex and subject to intellectual property rights, privacy regulations and other legal considerations.⁴⁸ Additionally, the cost of building or licensing a dataset from scratch can be prohibitive, making it difficult for smaller companies to compete with established players.

Moreover, there may also be antitrust concerns if the dominant players in the market control access to the datasets needed to develop AI models, as this could potentially stifle innovation and competition. Therefore, ensuring fair and open access to training data is a critical legal issue in the development and deployment of AI technology.

Another issue with input data is that while some large datasets are merely informational and not protectable, the majority of training datasets consist of copyrighted works. For instance, the corpus of works used to develop AI algorithms for text, facial recognition and image recognition all include copyrighted works. Thus, the question arises as to whether using these works is lawful and under what circumstances.

Currently, data collection for TDM has been considered fair use in the USA,⁴⁹ and there are exceptions and limitations under EU copyright law.⁵⁰ Specifically, in the USA, Google Books was granted permission to search entire libraries to provide search functions and

⁴⁴ I Goodfellow et al, *Deep Learning* (Cambridge, MA, MIT Press 2016).

⁴⁵ D Foster, *Generative Deep Learning: Teaching Machines to Paint, Write, Compose and Play* (Newton, MA, O'Reilly 2019) p 1 (stating that "a generative model can be broadly defined as a generative model describing how a dataset is generated, in terms of a probabilistic model. By sampling from this model, we are able to generate new data").

⁴⁶ *ibid.*, at 66 (illustrating how some companies have an advantage in the AI space because they are capable of using larger sets of training data to improve their algorithms, resulting in better performance. This gives them a "privileged zone" compared to other companies).

⁴⁷ *ibid.*

⁴⁸ See, eg, J Vesala, "Developing Artificial Intelligence-Based Content Creation: Are EU Copyright and Antitrust Law Fit for Purpose?" (2023) 54 *International Review of Intellectual Property and Competition Law* 351.

⁴⁹ See MW Carroll, "Copyright and the Progress of Science: Why Text and Data Mining Is Lawful" (2019) 53 *UC Davis Law Review* 893, 894 (arguing that fair use allows a TDM researcher to create non-transitory copies during processing and to preserve the processed data for archival purposes due to the transformative and beneficial nature of TDM). See also Lemley and Casey, *supra*, note 38, 746 (questioning whether machine copying will continue to be treated as fair use).

⁵⁰ See Art 4 of 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, *Official Journal of the European Communities* 2019 L 130, 92.

excerpts from books.⁵¹ However, it is unclear whether these conclusions apply to data collection and input for machine learning, as there is no copyrightable output. Indeed, it cannot be guaranteed that courts will apply this precedent to comparable technologies.⁵² In the USA, data collection for TDM may be permissible if it is a transformative use,⁵³ but it is not immediately clear that a copyrighted work is being transformed into another copyrighted work. In addition, in the Google Books case, the court recognised that Google's digitisation of copyrighted books, undertaken for the purpose of creating an extensive index and facilitating search functionality, constituted fair use.⁵⁴ This digitisation process was specifically designed to enhance users' ability to locate and access copyright owners' books, providing an invaluable tool for researchers, scholars and the general public. In this context, it is important to note that Google did not intend to compete with or replace the original works, but rather to improve their discoverability and enable consumers to make informed decisions about purchasing or accessing the entire works. On the other hand, when we examine generative AI technology, we encounter a contrasting scenario. Generative AI systems have the potential to empower users to easily produce content that may directly compete with the original ingested material. These systems utilise algorithms and machine learning techniques to generate new works, such as texts, images or music, based on the patterns and information gathered from existing content. Unlike Google's indexing and search functionality, which primarily served as a tool for information retrieval, generative AI opens the door for the creation of derivative works that could potentially overshadow or undermine the market for the original content. So, while Google's efforts in the Google Books case were found to align with fair use principles, the ease and accessibility of generative AI introduce complexities and challenges regarding copyright protection. Exactly for this reason, numerous court cases are currently underway in the USA seeking to clarify the definition of a "derivative work" and "transformative use" under intellectual property law, particularly with respect to copyrighted material used to train AI systems.⁵⁵ In particular, OpenAI and other prominent generative AI platforms are currently facing lawsuits alleging copyright infringement for training AI systems with illegally acquired datasets.⁵⁶ Specifically, in the legal case of *Tremblay v. OpenAI Inc.*,⁵⁷ the plaintiffs assert that OpenAI employed their copyrighted books without obtaining proper authorisation in order to train ChatGPT. The assertion is made that ChatGPT possesses the ability to effectively condense the content of various books, thereby implying that the chatbot has comprehensively engaged with and assimilated the information contained within said literary works. In the case of *Silverman et al. v. OpenAI Inc.*, the plaintiffs assert that OpenAI engaged in unauthorised utilisation of copyrighted work, specifically the book titled *The Bedwetter*, for the purpose of training ChatGPT.⁵⁸ Specifically, the authors of this class action claim that ChatGPT is capable of producing summaries of their novels when provided with a suitable prompt. They base this

⁵¹ See *Authors Guild v. Google, Inc.*, 804 F.3d 202, 214–15 (2d. Cir. 2015) (Google was authorised by the court to digitise all books available on the market, which served as an initial step towards creating a book search system that could provide exact excerpts of copyrighted text to users).

⁵² Lemley and Casey, *supra*, note 38, 763.

⁵³ A transformative use is one that "alter[s] the first [work] with new expression, meaning, or message". See *Campbell v. Acuff-Rose Music, Inc.*, 510 U.S. 569, 579 (1994).

⁵⁴ See *Authors Guild*, 770 F. Supp. 2d at 207–08.

⁵⁵ See, eg, *Getty Images (US), Inc. v. Stability AI, Inc.*, No. 1:23-cv-00135-GBW (D. Del. Mar. 29, 2023); *Silverman et al. v. OpenAI, Inc. et al.*, No. 4:23-cv-03416 (N.D. Cal. Jul. 7, 2023); *Tremblay et al. v. OpenAI, Inc. et al.*, No. 4:2023-cv- 03223 (N.D. Cal. Jul. 7, 2023). It is expected that the resolution of these litigations will depend on the interpretation of the fair use doctrine.

⁵⁶ *ibid.*

⁵⁷ See *Tremblay et al. v. OpenAI, Inc. et al.*, No. 4:2023-cv- 03223 (N.D. Cal. Jul. 7, 2023).

⁵⁸ See *Silverman et al. v. OpenAI, Inc. et al.*, No. 4:23-cv-03416 (N.D. Cal. Jul. 7, 2023).

claim on the fact that the AI tool has been trained using their copyrighted material, thereby establishing its familiarity with the content. Finally, in the dispute *Getty Images Inc. v. Stability AI*, the famous photo agency alleges that the software developer responsible for the AI art tool known as Stable Diffusion engaged in the unauthorised scraping of a substantial number of its images.⁵⁹ This act was purportedly carried out for the purpose of training the aforementioned system without obtaining proper permission or providing compensation to Getty Images. In addition, the AI tool Stable Diffusion generated a modified rendition of Getty's watermark, with the purpose of promoting, facilitating or concealing the infringement of Getty Images' copyright. This action – according to the plaintiff – also constitutes a violation of the Digital Millennium Copyright Act (DMCA) regulations regarding copyright management information.⁶⁰ Getty Images has also filed a similar complaint in the UK, requesting the High Court of London to issue an injunction barring Stability AI from selling its AI image generation technology in the country.⁶¹

The resolutions of all of these cases remain pending, and the manner in which they will be resolved remains uncertain at present. Nevertheless, these cases mark the initial significant legal confrontations regarding the utilisation of AI in relation to copyright violation. If the plaintiffs achieve a favourable outcome, they have the potential to exert a substantial influence on the advancement of AI technology.

However, the US Supreme Court's recent ruling in a non-technological case has already raised concerns about potential adverse implications on the intellectual property rights of works generated by AI.⁶² This case seems to have shifted the focus of the transformative use assessment. The controversy pertains to a conflict concerning copyright infringement, specifically regarding the utilisation of a photograph featuring the musician Prince that was taken in 1981.⁶³ The photograph was subsequently incorporated by the artist Andy Warhol in a series of prints and illustrations without obtaining the photographer's authorisation.⁶⁴ The fair use doctrine was invoked by the Andy Warhol Foundation for the Visual Arts to justify the creation of derivative works. In this context, the US Supreme Court ruled that the Foundation lacked a fair use defence to license a derivative rendition of the photograph for commercial purposes.⁶⁵ This recent decision could potentially result in a significant restriction of the transformative use doctrine, given that the Supreme Court appears to have effectively limited its scope.⁶⁶ So, it will be interesting to see what happens when US courts have to use the rules set up in this case to judge the licensing of AI training input. In the event that a court determines that data ingestion – which involves acquiring unprocessed data from one or more sources and modifying them to render them appropriate for the purpose of training AI machines – constitutes an act of infringement, the entire AI system may encounter significant legal difficulties. In fact, the vast majority of data that generative AI systems have assimilated – including both textual and visual content – have been de facto obtained without the express authorisation of the rights

⁵⁹ See *Getty Images (US), Inc. v. Stability AI, Inc.*, No. 1:23-cv-00135-GBW (D. Del. Mar. 29, 2023).

⁶⁰ See Section 1202(b) of the Digital Millennium Copyright Act, 17 U.S.C.A. § 1202(b). Specifically, the plaintiffs argued that the defendant's actions violated the provisions of the Digital Millennium Copyright Act (DMCA) because it altered or removed copyright management information (CMI) embedded in the plaintiffs' images and instructed the AI system to exclude any CMI from its generated output.

⁶¹ S Tobin, "Getty asks London court to stop UK sales of Stability AI system" (*Reuters*, 1 June 2023) <<https://www.reuters.com/technology/getty-asks-london-court-stop-uk-sales-stability-ai-system-2023-06-01/>> (last accessed 1 August 2023).

⁶² See *Andy Warhol Foundation for the Visual Arts, Inc. v. Goldsmith*, 143 S.Ct. 1258 (2023).

⁶³ *ibid*, p 1.

⁶⁴ *ibid*, p 1.

⁶⁵ *ibid*, p 2.

⁶⁶ See W Patry, "Andy Warhol Foundation for the Visual Arts, Inc. v Goldsmith: did the U.S. Supreme Court tighten up fair use?" (2023) 18 *Journal of Intellectual Property Law & Practice* jpad060.

holders. Consequently, here the question at hand pertains to the potential copyright infringement that may arise from utilising copyrighted works as training data. Specifically, it is necessary to determine whether such usage automatically constitutes copyright infringement or whether the distinct purpose of training data sufficiently diverges from that of the original copyrighted works, thereby warranting a fair use defence.

In contrast to the USA, the EU adopts a protectionist stance and has established a degree of accountability for the utilisation of training data. Specifically, the Directive on Copyright in the Digital Single Market (CDSM Directive)⁶⁷ includes Article 4(1), which provides a broad exception for TDM. Under this provision, individuals such as commercial AI system developers and educators may make copies of works or databases for the purpose of extracting information from text and data. They may retain these copies for as long as they are needed for the AI training process.⁶⁸ However, rights holders have the option to exclude TDM exemptions from their contracts with miners (ie entities or individuals that engage in TDM activities) in order to safeguard their commercial interests.⁶⁹ This particular provision has met with considerable criticism for providing a copyright exception that is perceived as being too restrictive. In contrast to the traditional understanding of copyright, which generally focuses on the protection of original expression, this provision appears to include factual information and data, and this aspect has drawn much criticism.⁷⁰ However, the manner in which this opt-out option can be implemented and the extent to which AI developers will adhere to it are still to be determined.

An additional issue associated with data aggregation pertains to the implementation of EU data protection legislation.⁷¹ Indeed, the process of data aggregation is of paramount importance in the training and refinement of generative AI models. This entails the gathering and merging of substantial quantities of data from diverse origins to augment a model's proficiency and functionalities. The processing of personal data within the EU is subject to stringent requirements and limitations, as stipulated by the General Data Protection Regulation (GDPR).⁷² These challenges remain unexplored in both doctrine and policy and need to be further explored and resolved.

VI. Source attribution and other copyright challenges in language models

ChatGPT's input as a language model comes from a variety of sources, including books, essays, websites and social media posts. These sources may contain copyrighted works that are used to train the language processing algorithms in ChatGPT. Given the legal concerns

⁶⁷ 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, Official Journal of the European Communities 2019 L 130, 92.

⁶⁸ See Arts 4(1) and (2).

⁶⁹ See Art 4(3) providing that "[t]he exception or limitation provided for in paragraph 1 shall apply on condition that the use of works and other subject matter referred to in that paragraph has not been expressly reserved by their rightholders in an appropriate manner, such as machine-readable means in the case of content made publicly available online". See also Bonadio et al, *supra*, note 38, 54.

⁷⁰ See T Margoni and M Kretschmer, "A Deeper Look into the EU Text and Data Mining Exceptions: Harmonisation, Data Ownership, and the Future of Technology" (2022) 71 GRUR International 685. As to the role of the idea/expression dichotomy in the generative AI debate, see also Lemley and Casey, *supra*, note 38.

⁷¹ See generally P Hacker et al, "Regulating ChatGPT and Other Large Generative AI Models" in *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*, 1112–23 (New York, Association for Computing Machinery 2023) <<https://doi.org/10.1145/3593013.3594067>> (last accessed 1 august 2023).

⁷² Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation) (Apr. 27, 2016), Art 99(2), 2016 O.J. (L 119) <<https://eur-lex.europa.eu/eli/reg/2016/679/oj>> (last accessed 1 August 2023).

surrounding copyright and the use of training data for machine learning, it is likely that ChatGPT faces similar issues. The dilemma, as with other AI systems, is whether using copyrighted material to train ChatGPT's language processing algorithms is legal and under what conditions. Because ChatGPT and other generative AI systems rely heavily on large amounts of training data, which may include copyrighted works, this presents a significant legal hurdle. The input data utilised by ChatGPT are produced via a method referred to as "training". During the training phase, the model is presented with a vast corpus of textual data, which are employed to instruct the speech-processing algorithms. The corpus under consideration exhibits the capacity to encompass a diverse range of text-based sources, including but not limited to books, articles, websites, social media posts and analogous materials. The type of data employed to furnish instructions to ChatGPT is dependent on the particular task or use case for which the model has been trained. If ChatGPT is directed to address customer service inquiries in a particular language, the training data corpus employed could be sourced from transcriptions of customer conversations or online evaluations. To guarantee adherence to copyright regulations, it is important to acquire any information utilised for the dissemination of ChatGPT through lawful channels. This may involve obtaining permission to use copyrighted materials or accessing publicly available information. Under specific circumstances, fair use or other legal exemptions may be relevant. However, it is essential to note that this is a multifaceted and dynamic field of law that necessitates meticulous examination on a per-case basis.

It is relevant to bear in mind that the programmers accountable for the development and training of ChatGPT hold the responsibility for ensuring that the training data remain free from any copyright violations. The provision of a comprehensive list of data sources may not be practical; however, OpenAI could explore more transparent options for disclosing the origins of the training data it employs. This could involve specifying the sources utilised or outlining the methodologies employed to gather and evaluate the data. This could reduce concerns about potential copyright infringement and improve transparency in the creation of AI models.

An intelligent individual writing creative content must provide a list of sources to prove the validity of their work and to avoid plagiarism. Why are ChatGPT language models exempt from this requirement? The answer is quite evident: language models lack personal convictions and the capacity to generate authentic ideas. Their capabilities stem from extensive training on diverse data sources, enabling them to generate texts. Nonetheless, it is crucial for any writer, be it human or machine, to acknowledge and reference their sources appropriately. This practice not only ensures the accuracy and reliability of the text, but also prevents some cases of plagiarism. However, there are significant differences in the way human writing and language models (eg ChatGPT) go about proofreading and citing sources. Human writers often view the source as a moral responsibility and express their accountability. They are responsible for the truth of their claims, and citing sources is one way to support their evidence.

In contrast, the issue of citing sources is somewhat more complicated for language models such as ChatGPT. Because these models generate text by using patterns and structures from training data, they do not inherently "support" particular claims or ideas. Instead, their answers are formulated based on statistical probabilities and patterns in the data. Consequently, it should be noted that ChatGPT's responses are not always accurate or reflective of reality, even though they come from an extensive data corpus. Therefore, it is important to emphasise the importance of source citations for language models. Incorporating this practice would effectively maintain the accuracy and credibility of the text, curb the spread of misinformation and provide transparency regarding the legitimacy of data sources. Sometimes sources can be provided automatically by using training data or contextual cues in the input text. In essence, the issue of source citation is equally important for human authors and language models. Although the approach and

timing of source citation may be different for these two groups, the basic principles of maintaining accuracy and credibility and avoiding plagiarism remain unchanged. It is also worth noting that there are different perspectives that challenge copyright holders' concerns about the use of their intellectual property in generative AI systems. Differing views arise from the fact that developers prioritise data encapsulated in copyrighted works over actual expression.⁷³ From the developers' perspective, documents and creative works are fundamentally viewed as collections of textual content, visual elements or auditory components that serve as unprocessed inputs to computational goals. The main goal of their research is to use the above raw material to train and extend generative AI models. This, in turn, facilitates the development of novel content by leveraging patterns and insights gained from existing works.

On the other hand, copyright law focuses on protecting the unique manifestation of a creation, commonly referred to as "original expression". This refers to the distinctive and innovative approach authors use to convey their concepts or create visual representations, melodies or other forms of artistic expression. It should be noted that copyright law is not able to cover the basic data, facts and concepts contained in copyrighted materials.⁷⁴

Proponents of using copyrighted materials for the purpose of training generative AI systems contend that copyright laws do not protect the basic data and concepts, and therefore it should be considered acceptable to use such works for computational purposes. The argument is that the focus is not on reproducing the exact form of the source material, but on using the information and structures present in that material to achieve novel and inventive results.

VII. Exploring alternatives for safeguarding AI training data

As previously discussed, using data for training purposes involves assimilating information obtained from publicly accessible websites on the Internet, including texts, images and other content. This procedure involves reproducing the content, which may violate the exclusive right of reproduction protected by copyright law and jeopardise the rights of authors and performers.⁷⁵ In order to advance and deploy generative AI across multiple industries, it is crucial to improve the access to and use of training data in terms of transparency and fairness. This is because AI algorithms rely heavily on enormous amounts of data to acquire knowledge and make accurate predictions. The availability and accessibility of such data are crucial factors in determining the efficacy and performance of an AI system. Keeping this objective in mind, we have considered a number of strategies for achieving this objective.

Establishing explicit data-sharing agreements with data providers is an essential first step.⁷⁶ Data-sharing agreements can be used to address the complex issue of using

⁷³ Regarding the role of the idea/expression dichotomy in the debate surrounding generative AI, see Lemley and Casey, *supra*, note 38, 772–76. See also MA Lemley, "How Generative AI Turns Copyright Law on its Head" available at SSRN <<https://ssrn.com/abstract=4517702>> (last accessed 1 august 2023; arguing that our current basic copyright doctrines – the idea–expression dichotomy and the substantial similarity test for infringement – do not fit generative AI).

⁷⁴ Lemley and Casey, *supra*, note 38.

⁷⁵ A recent open letter to policymakers on AI demanded creative rights in AI proliferation from various organisations and businesses that collectively represent more than six million artists, producers, performers and publishers worldwide. See International Confederation of Societies of Authors and Composers, "Global Creators and Performers Demand Creative Rights in AI Proliferation – An Open Letter to policy makers on Artificial Intelligence" <<https://www.cisac.org/Newsroom/articles/global-creators-and-performers-demand-creative-rights-ai-proliferation>> (last accessed 1 august 2023).

⁷⁶ For a view regarding the current data policies in the EU, see M Leistner and L Antoine, "IPR and the Use of Open Data and Data Sharing Initiatives by Public and Private Actors" (*European Parliament*, May 2022)

protected content for AI training while ensuring compliance with copyright laws and protecting content owners' rights.⁷⁷ These agreements are essential for delineating the scope of data usage, establishing limitations, specifying required permissions and arranging the necessary licenses for using copyrighted material in AI training processes.⁷⁸ Data-sharing agreements also enable AI developers to establish a legally binding framework that regulates the access, utilisation and administration of protected content throughout the AI training process. The implementation of agreements can yield advantages in terms of establishing unambiguous provisions pertaining to the authorised usage of data and guaranteeing that such usage is confined to the domain agreed upon by both parties. In addition, they have the potential to establish criteria for identifying non-viable data and assessing the permissible utilisation of content in the context of AI training. This encompasses determinations such as the permissible categories of AI algorithms or models and the criteria that govern the timing and duration of data utilisation. Furthermore, these agreements may comprise clauses that pertain to limitations on data utilisation and guarantee adherence to limitations enforced by proprietors of content by AI developers. These limitations may entail forbidding the retrieval or repurposing of data beyond the initial AI instruction or the dissemination or monetisation of protected material.

An additional viable measure for ensuring the protection of AI training data is to contemplate the implementation of certain types of remuneration programmes, such as revenue sharing or royalty payments, to guarantee that creators of copyrighted materials utilised in AI systems are duly compensated.⁷⁹ This strategy is important to demonstrate recognition of the inherent value of copyrighted content and to ensure that content creators receive a fair share of the benefits resulting from the use of their works by AI systems. AI developers can establish a direct correlation between the financial gains generated by AI systems and the use of copyrighted works by implementing revenue-sharing or royalty structures. The aforementioned scenario presents a persuasive motivation for content producers to provide their works as training data, given that they stand to gain directly from the financial prosperity of AI systems that utilise their creative output. In accordance with a revenue-sharing arrangement, creators of content would be entitled to a pre-established portion of the revenue produced by an AI system that is utilising their copyrighted materials. This could be a proportional arrangement in which the content creator receives a fair share of the revenue in proportion to their contribution to the training data. Such an arrangement ensures that content creators receive fair

<[https://www.europarl.europa.eu/RegData/etudes/STUD/2022/732266/IPOL_STU\(2022\)732266_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/STUD/2022/732266/IPOL_STU(2022)732266_EN.pdf)> (last accessed 1 August 2023).

⁷⁷ On the importance of data sharing, see "Towards a European Strategy on Business-to-Government Data Sharing for the Public Interest: Final Report Prepared by the High-Level Expert Group on Business-to-Government Data Sharing", COM (2020).

⁷⁸ M Kop, "The right to process data for machine learning purposes in the EU" (2021) 34 *Harvard Journal of Law & Technology* 1 (2021) (supporting the proposal for a right to process data for machine learning purposes).

⁷⁹ See M Sentfleben, "Generative AI and Author Remuneration" (14 June 2023) <<https://ssrn.com/abstract=4478370>> (last accessed 1 August 2023; proposing to introduce remuneration mechanisms that ensure the payment of compensation for the use of generative AI systems in the literary and artistic field); M Sentfleben, "A Tax on Machines for the Purpose of Giving a Bounty to the Dethroned Human Author – Towards an AI Levy for the Substitution of Human Literary and Artistic Works" (28 January 2022) <<https://ssrn.com/abstract=4123309>> (last accessed 1 August 2023); Sobel, *supra*, note 38 (supporting the idea of a scheme for authors who do not object to their works being used to train AI, but who want to be compensated); G Frosio, "Should We Ban Generative AI, Incentivise It or Make It a Medium for Inclusive Creativity?" in E Bonadio and C Sganga (eds), *A Research Agenda for EU Copyright Law* (Cheltenham, Edward Elgar, forthcoming) available at SSRN <<https://ssrn.com/abstract=4527461>> (last accessed 1 August 2023; exploring alternative mechanisms to support and promote human creativity in the face of AI advancements).

compensation for the value that their copyrighted works have on the functionality and success of the AI system.

Alternatively, a royalty-based compensation model could be implemented in which content creators receive a set fee for each use of their copyrighted works by the AI system. This fee structure could consist of a fixed amount per use or a percentage of the revenue generated by the AI system. This model guarantees that content creators receive fair compensation for the duration of the AI system's use of their copyrighted works by linking the fee to their usage.

The implementation of revenue-sharing or royalty structures requires explicit agreements between AI developers and content creators, specifying the exact terms of compensation. It is obviously important that these agreements specify the exact method for calculating revenue sharing or royalties, as well as their periodicity and temporal scope. Implementing transparent and mutually agreed-upon remuneration mechanisms can safeguard the interests of both AI developers and content creators, promoting a fair and sustainable ecosystem for integrating copyrighted works into AI systems. In general, the concept of compensating content producers through revenue sharing or royalties aims to recognise the importance of their copyrighted material in the AI ecosystem and to ensure that they receive adequate compensation for their role in the prosperity of generative AI tools. The scheme fosters a symbiotic connection between AI developers and content creators while maintaining equitable and just practices in the use of copyrighted material in the AI field.

An additional crucial policy element for safeguarding and reinforcing AI training data could potentially involve the creation and maintenance of open-source datasets intended for the purpose of training machine learning models. The majority of AI research is currently being funded by larger corporations. Hence, it is imperative to institute a programme that provides unrestricted or unobstructed entry to AI training data with the aim of fostering ingenuity, promoting collaboration and propelling the field of AI on a more democratic, equitable and transparent trajectory. The provision of such datasets could ensure that scholars, programmers and corporations are able to utilise them to construct and improve AI models. Advocating for the free or open accessibility of AI training data aligns with the principles of knowledge dissemination, accountability and equitable opportunities for both incumbent enterprises and emerging players from a policy perspective.⁸⁰ The reason for this is that well-established enterprises, owing to their prevailing market positions, possess extensive repositories of linguistic and visual data that can be leveraged for the advancement of AI. Enabling broad access to data has the potential to promote the progress of AI technology for the collective benefit of society rather than confining its advantages to a select few. This approach can also foster equitable competition by reducing entry barriers and facilitating the participation of smaller entities and marginalised communities in AI research.⁸¹ Hence, it is desirable to develop legislative measures that facilitate the promotion and the exchange of data while safeguarding privacy and intellectual property rights, thereby facilitating the utilisation of open-access AI training data. Specifically, these training data should be considered as a public “participatory good” because their production is based on collective efforts and their value results from the collective participation of numerous individuals who offered their creative content for the creation of training datasets.

⁸⁰ See OECD, “Recommendation of the council on artificial intelligence”, OECD/LEGAL/0449 <https://legalinstruments-oecd-org.sare.upf.edu/en/instruments/OECD-LEGAL-0449> (last accessed 1 August 2023).

⁸¹ See, eg, MA Lemley and A McCreary, “Exit Strategy” (2021) 101 Boston University Law Review 1, 68 (suggesting that specific AI training databases should be made accessible to all AI systems, or alternatively companies should permit their competitors to access these databases to ensure compatibility with the widely accepted standard).

To ensure the adequacy of data diversity for the training of AI models, it may be advisable for the law to incorporate provisions that incentivise businesses to voluntarily furnish anonymised data to publicly accessible repositories. The proposed framework has the potential to establish benchmarks for the ethical handling of data, safeguarding the privacy of individuals and preventing the deployment of AI applications that may lead to discrimination or harm. By establishing legal frameworks, it is possible to address issues related to data ownership and licensing, as well as defining the rights and responsibilities of both data providers and consumers. The implementation of similar legislative measures has the potential to foster collaboration among the government, industry and academia by means of funding schemes and recognition systems that prioritise initiatives pertaining to open-access AI training data. One potential strategy by which policymakers could foster a culture of openness and collaboration within the AI industry would be to provide grants and other incentives to researchers and organisations that prioritise the sharing of data. Under this perspective, a practical and concrete solution that could address the problem of copyright clearance of input data (training data) and give AI developers some breathing space would be to establish data repositories or clearinghouses for machine learning training datasets.⁸² Establishing data repositories or clearinghouses has the potential to make obtaining licenses and approvals much easier while promoting a more efficient and open process. Indeed, these repositories could act as centralised platforms that facilitate the process of obtaining licenses and permissions and enable negotiations between AI developers and content creators to take place more easily. They also play a crucial role in streamlining the complicated process of resolving copyright disputes, ensuring fair compensation and protecting the interests of all parties involved. Content creators have the option to formally register their works in a designated repository, where they must explicitly state the terms of use and compensation they expect to receive for the use of their protected intellectual property. This allows AI developers to easily access these data and ensure that negotiations are based on accurate and transparent information. AI developers and content creators can more effectively manage the complexity of compensation and rights issues through the use of data repositories. The centralised nature of these repositories promotes consistency and fairness in determining compensation and ensures compliance with licensing terms and copyright laws. In addition to the benefits already mentioned, data repositories can also promote fair competition among AI participants and provide good opportunities for new entrants in the field. Namely, these repositories can provide access to valuable datasets that smaller organisations might not have been able to access on their own. This creates new opportunities for innovation and competition, as a broader range of AI model and algorithm developers can leverage high-quality data, reducing the concentration of data ownership and promoting competition in the AI industry.⁸³

⁸² See, eg, M Kop, “Machine Learning & EU Data Sharing Practices” (3 March 2020) Stanford–Vienna Transatlantic Technology Law Forum, Transatlantic Antitrust and IPR Developments, Stanford University, Issue No. 1/2020 <<https://ssrn.com/abstract=3409712>> (last accessed 1 August 2023; briefly discussing the establishment of an online clearinghouse for machine learning training datasets).

⁸³ A concrete example of initiatives in the field of open data and data sharing for AI training data is the Global Initiative on AI and Data Commons. See International Telecommunication Union, “Global Initiative on AI and Data Commons” <<https://www.itu.int/en/ITU-T/extcoop/ai-data-commons/Pages/default.aspx>> (last accessed 1 August 2023). See also The European AI Alliance, European Commission <<https://futurium.ec.europa.eu/en/european-ai-alliance>> (last accessed 1 August 2023).

Ultimately, implementing ethical guidelines and industry standards for AI training could serve as an additional viable means of protecting AI training data.⁸⁴ These guidelines address issues such as data collection, consent, fair compensation and the transparency of the training process. Adherence to these guidelines can help to ensure responsible and legal use of copyrighted works, builds trust and foster a culture of compliance and respect for copyright.

Consent is another essential element that can be covered by such ethical principles. Legal use of copyrighted works requires obtaining appropriate consent from content creators or affected individuals. Such ethical guidelines provide guidance on obtaining informed consent that ensures content creators fully understand how their works will be used in AI training processes and allows them to give consent or opt out if they wish. This promotes integrity and respect for the rights of content creators and fosters a mutually beneficial relationship between AI developers and creators of copyrighted works.

Another fundamental aspect addressed in such ethical guidelines is transparency. These guidelines would promote transparency in AI education by ensuring that the sources and uses of copyrighted works are properly documented and disclosed. Transparency increases accountability and trust among stakeholders and allows content creators and AI system users to track the use of copyrighted works and resolve potential concerns or disputes.

Despite significant advances in the field of generative AI, there is still a reluctance to deploy it widely. A significant obstacle that requires attention is the lack of clearly defined standards in critical areas, including but not limited to intellectual property rights. Successful management of copyright protection for training data can be facilitated through the implementation of effective policies by policymakers. This can ensure the smooth incorporation of generative AI systems. The aforementioned proposals seek to highlight the importance of clear regulations in this context. In particular, these proposals could help to enable fair and clear data retrieval, equitable remuneration of creators, simplified licensing procedures and compliance with copyright regulations and ethical principles. When formulating AI regulations and policies, policymakers should consider the above factors.⁸⁵ By incorporating these variables into their models, policymakers can create a durable and innovation-driven environment for the advancement of generative AI.⁸⁶

⁸⁴ This particular aspect has attracted significant attention both on a global scale and within the EU. See, eg, OECD, Observatory of Public Sector Information, “Algorithmic Impact Assessment” (2019) <oecd-opsi.org/toolkits/algorithmic-impact-assessment/> (last accessed 1 August 2023); EC, “EU Guidelines on ethics in artificial intelligence: Context and implementation” (2019), PE 640.163.

⁸⁵ The current discourse on AI governance seems to largely overlook these particular aspects. The EU’s AI Draft Regulation (see *infra*, note 86) does not discuss intellectual property issues in depth, as the current version contains only a general obligation to disclose summaries of the copyrighted information used for training purposes. The USA has taken steps to address the risks associated with AI, with the National Institute of Standards and Technology – within the US Department of Commerce – developing a voluntary AI risk management framework. This is the first explicit US government guideline on standards in AI system design, development, deployment and use. See Department of Commerce’s National Institute of Standards and Technology, Artificial Intelligence Risk Management Framework (AI RMF 1.0) (January 2023) <<https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>> (last accessed 1 August 2023). In addition, the White House Office of Science and Technology Policy has published non-binding guidance for an AI Bill of Rights. See “Blueprint for an AI Bill of Rights”, Office of Science and Technology Policy <<https://www.whitehouse.gov/ostp/ai-bill-of-rights/>> (last accessed 1 August 2023). For a more comprehensive summary of regulatory policies and actions linked to AI in the USA, see CS Yoo and A Lai, “Regulation of Algorithmic Tools in the United States” (2020) 13 *Journal of Law and Economic Regulation* 7, 7–9.

⁸⁶ For a first concrete example of a piece of law designed to mitigate the potential hazards and challenges associated with the advancement and execution of AI, see Proposal for a Regulation of the European Parliament

VIII. Conclusions

This article examined copyright issues related to generative AI from the general perspective of the ChatGPT case study. It presents methods for addressing legal challenges in the development of AI systems, with the goal of protecting both copyright holders and competitors. The first part of the paper explored both the theoretical and practical applications of complex language models such as ChatGPT. The second part looked at the output of the ChatGPT model and discussed copyright issues. The third part looked at the training data and discussed copyright concerns.

We have also emphasised the increasing number of legal actions targeting generative AI systems. The litigation in question specifically focuses on the developers responsible for creating these systems, including ChatGPT. A significant number of cases involve various aspects of copyright protection, such as the training data used to train AI models and the nature of the data employed for this purpose. The research concludes that the ethical and legal concerns raised by AI model development must be addressed holistically, considering both inputs and outputs. The management of intellectual property rights in AI goes beyond outputs to include inputs, namely training data. AI systems rely heavily on large amounts of training data, which often include copyrighted works. This raises questions about the lawful collection and use of such data, as well as the creation of derivative works during the training process. Access to large datasets has become a competitive advantage for incumbents that can hinder innovation and competition. Ensuring fair and open access to training data is critical to the development and deployment of AI technology. The creation of AI training data is indeed a collective effort that requires the participation of many individuals. The value of these data comes from the collective participation of many individuals who have provided their creative content for the creation of training datasets, and they should therefore be considered shared resources that are available to all. Their use should be guided by principles of fairness and transparency.

Current legal frameworks, such as fair use in the USA and the TDM exemption in the EU, provide some guidance on the use of copyrighted material to train AI models. However, these frameworks may not fully address the complexities inherent in generative AI systems, which can directly compete with and even dwarf original works. Balancing technological advances with the preservation of creators' rights is critical to navigating the copyright landscape in the context of AI. In particular, finding alternatives for protecting AI training data is critical to improving transparency and fairness in data access and use. Strategies such as clear data-sharing agreements, compensation models that provide for revenue sharing or royalties, data repositories or clearinghouses and the development of ethical guidelines and industry standards can promote responsible and lawful use of copyrighted works in AI systems. These approaches ensure compliance with copyright laws, protect the rights of content creators, streamline licensing procedures and promote a sustainable and innovation-friendly AI ecosystem.

It is becoming increasingly clear that the growing capabilities of machine learning systems raise concerns about potential copyright restrictions. As AI technology continues to advance, the use of copyrighted works in training data is becoming more common, leading to the need for robust mechanisms to protect intellectual property rights. In the future, it will also be necessary to emphasise the responsibility of AI developers to be

and of the Council: Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts, European Commission (22 April 2021), 2021/0106(COD) <eur-lex.europa.eu/procedure/FI/2021_106> (last accessed 1 August 2023) [EU AI Draft Regulations] (which puts obligations on providers and distributors). For some comments from a comparative perspective, see F Patel and I Dyson, "The Perils and Promise of AI Regulation" (*Just Security*, 26 July 2023) <<https://www.justsecurity.org/87344/the-perils-and-promise-of-ai-regulation/>> (last accessed 1 August 2023).

proactive with their data-sourcing methods. It will be important for AI developers to implement methods to ensure the provenance of AI-generated content in order to provide more clarity about the works contained in training data. In the face of a new technological dilemma, copyright once again has a critical role to play in reconciling the competing interests of content creators and AI developers. This involves protecting the integrity of original works, ensuring adequate remuneration and addressing the potential dangers and complexities associated with the rapid advancement of generative AI. Consequently, such technology will require a profound paradigm shift in our conception of creativity and a corresponding re-evaluation of our approach to copyright.

Competing interests. The author declares none.