
Miscellaneous

Bruno Caldas Vianna

<https://orcid.org/0000-0003-0213-6115>

bruno.caldas@citm.upc.edu

Universitat Politècnica de Catalunya

Carles Sora-Domenjó

<https://orcid.org/0000-0003-2761-2384>

carles.sora@citm.upc.edu

Universitat Politècnica de Catalunya

Josep Torelló

<https://orcid.org/0000-0001-5662-4806>

josep.torello@citm.upc.edu

Universitat Politècnica de Catalunya

Submitted

May 31, 2024

Approved

July 4, 2025

© 2026

Communication & Society

ISSN 0214-0039

E ISSN 2386-7876

www.communication-society.com

2026 – Vol. 39 (1)

pp. 191-204

How to cite this article:

Caldas Vianna, B., Sora-Domenjó, C., & Torelló, J. (2026). Ethical and societal challenges in the use of generative AI in communication, *Communication & Society*, 39(1), 191-204.

<https://doi.org/10.15581/003.39.1.014>

Ethical and societal challenges in the use of generative AI in communication

Abstract

This paper presents a narrative review describing the challenges of adopting artificial intelligence tools for audiovisual communication, identifying issues related to the ethics, control, transparency and legal framework. We make a critique of the term and delimit the research to generative neural-network-based computational processes. The problems are examined in three proposed areas: the input, which includes the training material; the process, which deals with the systems necessary to run these services, their development and control; and the output, in which the use of generated artifacts is analysed. International legal and judiciary approaches are reviewed in each step, as they are fundamental for the development of the study. Finally, some hypothetical future scenarios based on current AI debates are imagined, which help readers foresee possible situations regarding the use of generative systems. We conclude that although it is unstoppable that these technologies are adopted, their use might have some limitations, regulations or conditions. More importantly, the general use of AI tools will bring a shift in the way we produce and consume information. While we can clearly foresee dramatic increases in productivity, it is not clear how this production will be accepted by society. On one hand, studies have already revealed that subjects rank AI-generated content higher than content made by humans. On the other hand, the plethora of artificial media might provoke an overall rejection of the digital world.

Keywords

Artificial intelligence, ethics, society, communication, neural networks, copyright.

1. Introduction

The goal of this narrative review article is to understand the legal, ethical and societal challenges involved in adopting media tools based on artificial intelligence in particular generative procedures. The narrative review methodology was chosen as it is “often useful for topics that require a meaningful synthesis of research evidence that may be complex or broad and that require detailed, nuanced description and interpretation (Sukhera, 2022)”.

Artificial intelligence systems that create content have been widely and rapidly adopted since the invention of Generative Adversarial Networks in 2014 (Goodfellow et al.,

2014). While originally limited to image production, neural networks have been adapted to produce music (Mittal et al., 2021), voice (Liu et al., 2022), texts (Floridi & Chiriatti, 2020), and videos (Blattmann et al., 2023). Moreover, the invention of multimodal models that encode textual references into different representations, such as CLIP in the case of images (Radford et al., 2021), has enabled the emergence of methods that use prompts, which are written descriptions used to guide the generation towards the desired result. This has triggered an explosion of commercial services aimed at the professional application of generative AI in the audiovisual industry.

Before proceeding with this analysis, it is necessary to briefly define the term artificial intelligence. By defining this set of computational processes as an artificial alternative to human cognitive capabilities, the field has opened the doors to endless comparisons between machine and human tasks. This, in turn, has led to inconvenient anthropomorphic analogies (Watson, 2019). However, the reasoning processes of computers and biological brains, despite the initial neuronal inspiration, are completely different, and they would benefit from being analysed as two separate domains. Likewise, it would be beneficial for this conversation to adopt a different term, like computational cognition or even the continuity of cybernetics, which is now seen as a mere antecedent of AI. Although it will be used throughout the text, we acknowledge the problems with the term artificial intelligence.

Another necessary distinction is that although there is a vast number of machine learning applications, they can be superficially distinguished into two types. The first, predictive AI, uses massive amounts of data to interpret the world, and can be found in applications such as data analysis and classification, computer vision and speech recognition. The second, generative AI, is used to produce new artifacts, also based on massive data, and will be the focus of the article.

To tackle the complexity of generative audiovisual production and its current debates, we analyse the challenges by dividing them into three types: input, processing and output. The input consists of the data used to train the models. Processing refers to the systems used for generation, including software, the neural architecture and the weights, that is, the neural database that results from the training process. Finally, the output is the resulting artifacts: images, videos, sounds, texts. We look at how their use and access could become limited. Entangled issues of ethics, copyright and legislation appear for all of these. Our objective is to pull from these threads and propose some future scenarios of how they can be dealt with.

2. Input

We describe here the process of feeding data to build the models used in audiovisual creation (Bommasani et al., 2022). Generative artificial intelligence also has a secondary input to the creative process: the parameters, including the prompts, that are used in the inference phase, when images and text are generated. However, this happens only once the models have been trained, and we will look at this later in this text.

To work properly, the systems need to be based on millions of existing examples of images, songs, texts and more. In fact, one of the factors that explains the AI boom is exactly the availability of this kind of data on the Internet. The other factors include the popularization of graphical processing units, GPUs and the innovative methods developed in the 2010s (Dean, 2022; Kautz, 2022). However, while this data might be easily available on the web, using it is not necessarily free. The fact that a given material is published does

not mean that it can be used for any purpose. The 1886 Berne convention, ratified by a vast majority of countries, established that works are protected from the moment of creation, no matter if they have been registered officially or not (*Berne Convention for the Protection of Literary and Artistic Works*, 1886). This means that unless the author states clearly at the moment of publication that the creation is available under specified conditions (usually defined in permissive licences such as Creative Commons) the material cannot be used for anything else. The material legally available is described by scholar Antonios Broumas as the intellectual commons (Broumas, 2020), and in some cases, such as with programming code and text (in Wikipedia articles, for example), it is so vast that it can be enough data to train certain applications.

However, there has been a strong reaction from artists, such as illustrators, against the appropriation of unlicensed productions for use in training AI models (Jiang et al., 2023). At the moment of writing, at least in the United States, it seems like the outcome of this conflict will be determined by the decisions of several open litigations, such as the one involving artists Karla Ortiz, Grzegorz Rutkowski and others against companies such as Stability AI and Midjourney¹. If there is a ruling in favour of the litigators, this could hinder the training methods used by these companies. Models that have already been trained, which can be found and downloaded from several websites and exist in the computers of countless users, could become a bootleg item, and using them could be considered illegal.

The other side to this scenario is the licensing of content to be used for training, which seems to be the outcome desired by the litigators. There are different approaches to this strategy. In the case of image generation, companies like Adobe already have the rights to vast repositories of intellectual property, which, added to the intellectual commons pool, is enough for developing proprietary models. This is the case of products and services such as Firefly and the generative tools recently incorporated into Photoshop (*Adobe Unveils Future of Creative Cloud with Generative AI as a Creative Co-Pilot in Photoshop*, 2023).

This alternative has its own set of limitations. Since the training set only contains the intellectual property of Adobe, no other visuals registered to other copyright owners can be generated. Firefly won't produce images of Nintendo's Mario, or Disney's Donald Duck. While this has the positive effect of limiting the abusive use of the tool in creating copyright-protected images, there are cases in which these creations are acceptable, as for instance in satire and critique. A cartoonist who wants to mock Homer Simpson's sexist attitudes would have to resort to a different tool. Likewise, even characters whose rights have already expired, such as Mickey Mouse, cannot be recreated in their platform, at least at the moment.

Another option is available for companies such as OpenAI or Google, which have virtually unlimited funds and can afford to license materials from media conglomerates. While the New York Times is suing OpenAI for using their articles (Grynbaum & Mac, 2023), Apple is reported to be in talks with major publishers such as Condé Nast and NBC News with the aim of training their own models (Mullin & Mickle, 2023). In February 2024, Google reached a deal with on-line forum Reddit to license the posts of their users (Tong et al., 2024). The social network's filing for their initial public offering of stocks² (IPO)

¹ <https://imagegeneratorlitigation.com/>

² <https://www.sec.gov/Archives/edgar/data/1713445/000162828024006294/reddits-1q423.htm>

confirms that it has already made 203 million dollars from such licensing agreements. One of Elon Musk's motivations for acquiring the social network Twitter (now renamed *X*) was certainly his desire to access its content to use it for training his own artificial intelligence startup, *xAI*.

This concept of using social media posts to train neural networks is a perfect example for demonstrating that legal limitations are not enough to dismiss ethical concerns (Coeckelbergh, 2020, Chapter 7). Even if the terms of use of sites like Facebook, X and Instagram explicitly allow the licensing of materials published in their platforms, many users express discontent with the fact that their cognitive output is turned into a profit that exclusively benefits the shareholders of these corporations. It is a textbook example of knowledge labour exploitation (Fuchs, 2010).

It is important to note that the use of unlicensed material is not necessarily illegal, which is actually the reason behind so many legal actions. Human cognitive production has never been used in this way or on this scale. Copyright law could not foresee this use, which could be considered *fair use*. This legal concept³ deals with uses which do not jeopardize the commercial activity of copyright holders, or which are otherwise useful for society, including, for instance, educational purposes or news reporting. Fair use is not defined in law, but rather decided case by case. Hence, the need for jurisprudence on these novel applications, as in the case of AI, which, given the scope of its consequences, could reach the Supreme Court of the United States.

It is noteworthy that the situations described so far use the United States legal framework as the reference. Other jurisdictions, however, offer different views of the problem. One of the clearest examples is the European Union directive on Copyright on the Digital Single Market (CDSM)⁴, which came into force in 2019. The directive defines the broad category of Text and Data Mining (TDM) as “automated computational analysis of information in digital form.”

The text defines two clear categories of exceptions to the reach of copyright in two of its articles. Article three states, in general terms, that academic institutions, as well as cultural heritage organizations, can perform TDM on copyrighted materials for research purposes. That was the path taken by the CompVis group at the Heidelberg University in Germany (currently at the Ludwig Maximilian University in Munich) to train the first text to image diffusion models (Rombach et al., 2022), later known as Stable Diffusion. Article four makes an exception for TDM on copyrighted assets even for commercial use as long as copyright owners have the possibility to “opt-out”, that is, to disallow this use of their creations. The text leaves out the details of how the opt-out mechanism would work; however, the company Spawning.ai⁵ started a database of artwork URLs whose owners do not want their content to be used in the AI training sets. Laion.ai⁶, a German non-profit that compiles web materials into datasets, which were the basis for the original diffusion models created by Rombach and his group, has incorporated Spawning's database, therefore removing opted-out images from some of their collections. The result is an AI model which respects the consent of creators while displaying greater aesthetic diversity than the ones trained on sourced content.

³ <https://www.copyright.gov/fair-use/>

⁴ <https://eur-lex.europa.eu/eli/dir/2019/790/oj>

⁵ <https://spawning.ai/>

⁶ <https://laion.ai/>

Japan sits at the other extreme in this regard. In 2019, a law was passed which allowed training on any copyrighted material, with a clear intention of making the country a safe haven for the development of AI. The law was clarified in 2023 in response to reactions from the content creator community, and included clauses that the training must not create unjust harm, conflicting “with the copyright holder’s market or inhibiting potential marketing channels for such copyrighted works in the future”. It also prohibited training with the “purpose of enjoyment of the thoughts and or sentiments expressed in the copyrighted work”, which might be difficult to interpret (Warren & Grasser, 2024).

In general terms, copyright only protects the expression of ideas, but not the concepts or ideas themselves. In other words, although the fixed representations are protected, the styles they convey are not (Wolfson, 2023). And this is at the core of the dispute: artists who have developed a particular style feel threatened because their concepts can be reproduced by simply using their names in the prompt. The aforementioned Greg Rutkowski found his name tagged 93,000 times in the training of Stable Diffusion (Heikkilä, 2022). On the other hand, sampling styles is exactly what allows art to evolve. In our view, prohibiting being able to reference styles developed by others would certainly stifle artistic innovation. Unfortunately, recent court decisions in the US have granted artists extensive protection over signature traits that represent an overreach of the original concept of copyright. The state of musician Marvin Gaye won a dispute with Pharrell Williams, claiming that although their song *Blurred Lines* had no melodic or harmonic plagiarism, it imitated Gaye’s signature style (Zernay, 2017). Being an unforeseen decision, the case set a legal precedent that allowed several other similar suits to follow (Elbeshbishi, 2022). Although it might protect the earnings of established artists, it could also threaten artistic innovation.

A significant societal downside to the adoption of AI tools is their environmental footprint. The training of models is highly computationally intensive, and therefore it is an energy hungry process. However, as we will see in the next section, AI companies are becoming increasingly less transparent about their training methods, making it difficult to properly assess their effect on the environment. The documentation for early models that is still publicly available shows that the training of OpenAI’s GPT-3 used a total of 1,287 MWh, the equivalent of three commercial jet round trips from New York to San Francisco (Patterson et al., 2021). This model has 175 billion *parameters* (the standard measure of model size), while GPT-4 is estimated, by independent sources, to have 1.7 trillion parameters (Schreiner, 2023), making it bigger by at least one order of magnitude. It is worth mentioning that new, smaller, models are becoming more efficient: Meta’s new Llama3 has quality versions with 8 and 70 billion parameters whose training, according to their *model card*, emitted the equivalent of 2290 tons of CO₂ (tCO₂eq)⁷, while GPT-3 used 552.1 tons, representing a four-fold increase.

A different kind of threat comes from the private, individual generation of data through AI processes. The ease with which useful content can be generated makes it very tempting to mass produce webpages and social media posts. Once they’re published, search engines track them, and they can be scraped for new trainings. The problem is that this kind of endogamic, recursive learning is known to cause models to “collapse” (Shumailov et al., 2024). In other words, the output of such training is useless garbage. As human-produced content becomes more difficult to filter out, we might get to a point

⁷ <https://huggingface.co/meta-llama/Meta-Llama-3-70B>

where the strategy of using online materials for training will no longer be viable. The other side to this threat is consumer burn-out itself, which will be described in the output section.

3. Process

The second group of threats to adopting AI creative processes is related to the control over the software and data used to run the models and the level of transparency they display. The first generation of neural models came either from academic research or from companies committed to an altruistic development of models. The founding declaration of OpenAI, from 2015, said that their “Researchers will be strongly encouraged to publish their work, whether as papers, blog posts, or code, and our patents (if any) will be shared with the world⁸.”

The harsh reality is that after the original Stable Diffusion from Heidelberg University and OpenAI’s GPT -2, few completely transparent models have been published. Not only does this mean that academics cannot study or replicate results obtained from these systems, but also that other developers cannot build new innovations on them. The main reason behind this is that generative AI has become immensely profitable, and there is fierce competition among the companies to deliver better products. However, another important factor is their unwillingness to share what is the underlying data powering the models. As we have seen, there is still no clear definition on the legality of the training sets. In this sense, companies prefer to avoid disclosing where their images, texts and other data are coming from. It is likely that only legal enforcement will be able to shed some light on the materials used.

To mitigate the negative perception caused by closing models, some companies choose the strategy of cloaking the whole process, but then publishing the file with its weights. This is the resulting “database”, which is the outcome of the training process. This allows end users to make inferences, that is, to use prompts to obtain outputs from the models—essays, texts, songs. The idea here is not only to offer some kind of free access, but also to allow independent tinkerers or even other companies to make improvements and adaptations on the models, which can in turn benefit the original publisher. This is clearly the case with the leak of the original Llama model developed by the company Meta (Hern, 2023). In some social media postings, its CEO Mark Zuckerberg praised the open-source approach of his corporation⁹. There is currently a plethora of large language models such as Phi, Mistral, and Llava, besides Meta’s Llama, which can be locally run on personal computers, given the user has enough processing power¹⁰. Some of these even come from other computer behemoths, such as Microsoft.

Another open-weight champion is the original enabler of Stable Diffusion, Stability AI. They not only provided the university researchers with the high-performance computing cluster needed to train version 1.5 of Stable Diffusion but also hired their top researchers to continue to develop their products in-house. Stability, like Meta, has released all their weights to the public since then, and the response was extraordinary: hundreds of their customizations can be found and downloaded from specialized sites such as Civitai.

8 <https://web.archive.org/web/20180328155139/https://blog.openai.com/introducing-openai/>

9 <https://www.instagram.com/reel/C2QARHJR1sZ/?stream=top>

10 <https://ollama.com/library>

The last problem of working with closed or semi-closed foundational models is that there is no way of knowing whether the training process was ethical, not only in terms of copyright issues. Many companies use underpaid workers to limit the output of disturbing content like “child sexual abuse, bestiality, murder, suicide, torture, self-harm, and incest” (Perrigo, 2023). Unless the methods and data used become transparent and well-documented, it is not clear whether these labour malpractices or environmentally unsound energy have been used, or even if financing has come from unlawful actors.

Completely closed models take away the users’ control in the sense that the users are dependent on the use that the company is offering. A system built on top of a given version of ChatGPT that, for instance, performs automatic translations, may stop working if the provider updates the models and discontinues the version that was being used. An artist may take months in refining their ability to generate images on a given version of Midjourney (a company that does not share their weights) only to see it removed at some point. Moreover, the open ecosystem advances so quickly that corporations are behind in incorporating most features into their commercial offerings. One example would be the posing extension *ControlNet*, a powerful open-source visual generative control tool that is being slowly incorporated into commercial systems.

A different procedural threat is related to the supposed existential risks of developing Artificial General Intelligence, that is, what could happen when AI’s capabilities reach human levels and systems can make decisions and control societal processes independently. Researchers Timnit Gebru and Émile P. Torres recently proposed the acronym TESCREAL, standing for ‘transhumanism, extropianism, singularitarianism, cosmism, rationalism, effective altruism, and longtermism’, for this set of beliefs that are popular among Silicon Valley leaders such as Ethereum’s Vitalik Buterin and Elon Musk (Gebru & Torres, 2024). According to this creed, the development of AI must be carefully and ethically guided to prevent it from becoming an existential threat to humanity. The creation of OpenAI has deep roots in these tenets, hence their original emphasis on openness. The 2023 ‘coup’ at the company, when CEO Sam Altman was temporarily ousted (Metz, 2023), can be attributed to the hard TESCREAL members of the board who believed that their product was being intensely commercialized with no thought for the risks involved. Believers of this theory support strict regulation of the development and use of AI.

We also need to evaluate the environmental impact of AI at the time of inference, that is, when the trained models are used to generate new content. OpenAI and Google, for instance, do these operations in huge data centres and are not transparent about their power consumption. Right now, AI use represents only 2% of global datacentre usage, but that number is expected to rapidly increase to 10% (Sisson, 2024). Moreover, these large servers use vast amounts of water to cool down, often in regions where water is already scarce. In 2027, global water use for AI datacentres could reach the equivalent of the water withdrawal of the country of Denmark (Ren, 2023).

As in the case of training, it should be noted that some models are becoming smaller, more open, and more efficient (like the aforementioned Llama-8b or 13b, or models from Mistral), so much so that in many cases AI users are moving from cloud providers towards self-hosted solutions in consumer PCs at home or work. These PCs are no different from computer gaming rigs in terms of power consumption. The effect of the shift remains to be studied: gaming implies a long continuous use, while AI is used in bursts as users need to ask questions or generate content. However, gaming is a leisure activity carried out in free time, while AI is used in the longer, working hours.

4. Output

The main threat to accessing AI-generated content here lies in the other end of the copyright diatribes: who can claim ownership to the materials created by neural networks? The question of the authorship of works partially made by machines is not trivial and has been pursued by different thinkers. Galanter, for example, proposed that in the case of deep learning, a good portion of authorship belongs to the system (Galanter, 2016). Zeilinger describes an arrangement where the artist responsible for the work *All We'd Ever Need Is One Another*, Adam Basanta, is a mere participant in the art making process (Zeilinger, 2021). In contrast, a recent doctoral thesis on visual arts and AI follows Ada Lovelace's suggestion that the *origination* is always human (Caldas Vianna, 2024).

Notwithstanding the concerns of academia, the commercialization of creative endeavours is a business in which authorship must always be clearly defined, as this will, in its turn, define ownership. And that is where the problems start.

Generative AI is a process in which not only countless humans take part, but also in which the systems make important contributions to the final result. On the human side, there are all the original artists, photographers, writers and social media dwellers whose materials were ingested in the training phase. Then we have the programmers who prepared the code that enables the system to run. Finally, there are the scientists who developed, on the shoulders of precedent scientists, the methods needed to mechanically spawn new cognitions and shared these methods through papers and journals. In the moment of inference, we have an immediate originator of the work: the amateur or specialist that refines the perfect prompts, sets parameters and also necessarily selects the most meaningful results from the endless possible outputs.

On the mechanical side of these processes, identifying the role of the machines is more complicated. While practitioners of computational generative art before AI always purposefully introduced elements of randomness to surprise both the audience and the creators themselves (Galanter, 2003), we cannot say that this gives computers agency. Actually, generating random numbers is not a trivial task for computers: they quite often rely on human or natural inputs to obtain some type of contingency. Works based on statistical models, on the other hand, are, by definition, unpredictable. In fact, both textual and visual generation tools allow the level of "unpredictableness" of the model to be controlled (CFG in image diffusion, temperature in LLMs).

The only reason we can still replicate generative AI processes and obtain exactly identical results is that the computation is done on binary, discrete machines. However, right now, there are optical chips in the works that perform the linear algebra needed for neural operations much more efficiently in light beams and sensors than in transistors (Brückerhoff-Plückelmann et al., 2023; Chen et al., 2023). When inference starts to be carried out in these chips, their computation will be intrinsically non-deterministic.

However, even in that case, it is difficult to argue that the machine is responsible for the output. A human-in-the-loop is the model that applies to generative AI, where a human has decided on a subject, written a prompt, and chosen a particular style by choosing the model used. And that is the reason it is difficult to accept the stance of the United States Copyright Office (USCO), when it affirms that machine-generated works aren't eligible for copyright protection (USCO, 2021).

The reference case that illustrates this limitation was the registration of the comic book *Zarya of the Dawn*, authored by Kristine Kashtanova. While the story was created by Kashtanova, all the illustrations were created using Midjourney as a generative AI tool.

The work was submitted for registration at USCO, without any indication of the tools used. Once the protection was granted, the author publicized the fact that the work was technically the first AI generated work to fall under the copyright domain. However, once the Office realised the repercussions of this, the registration was temporarily suspended for “further analysis”. Following a few months of deliberation, the decision taken in February 2023 was that the story and graphical composition of the comic was under protection; the images, however, were not (Lindberg, 2023). The arguments for their decision were related to randomness: “the process is not controlled by the user because it is not possible to predict what Midjourney will create ahead of time” or “Rather than a tool that Ms. Kashtanova controlled and guided to reach her desired image, Midjourney generates images in an unpredictable way.”

At the heart of the issue is their official take on what constitutes an author, which posits that “To qualify as a work of ‘authorship’ a work must be created by a human being” (USCO, 2021)—as if the machines were not created and operated by humans. Of course, not being able to register and therefore commercially exploit AI-generated works might have consequences for their adoption by the audiovisual industry. Professionals would be limited to using such creations as references and concepts, since there would be no guarantee that an animated cartoon exhibited on a streaming platform would have copyright protection. No Hollywood studio would undertake the production of features using OpenAI’s Sora video generator¹¹ if it couldn’t limit their reproduction. Creators would still have the option to build on these assets—much like Kashtanova developed her comic book—but would still run the risk of having essential works being appropriated commercially.

Nevertheless, as in the case of copyrights over training materials, different jurisdictions take different stances. The United Kingdom’s law says that “In the case of a literary, dramatic, musical or artistic work which is computer-generated, the author shall be taken to be the person by whom the arrangements necessary for the creation of the work are undertaken (Copyright, Designs and Patents Act, 1988)”, clearing the way to grant authorship to prompt artists. A survey of the Spanish jurisdiction, however, concluded that it would be difficult to define how much human input would be necessary to surpass the threshold required for a work to be protected (Calleja Reina, 2023).

Another challenge is the subjective acceptability of these products by society. Will viewers feel contempt for works which were visibly engendered by neural networks? Recently, world renowned graffiti artist Kobra posted a computer-generated image in his social media and suffered a dire backlash from his followers, leading to the later deletion of the post¹². But what happens when the generative origin of the product is undetectable, as in texts, voice overs or even realistic pictures? Would the evaluation of the work change when the audience is informed that it is generated? That was the exact subject of a recently published study, which found that “despite generally preferring the AI-generated artworks over human-made ones, the participants displayed a negative bias against AI-generated artworks when subjective perception of source attribution was considered, thus rating as less preferable the artworks perceived more as AI-generated, independently of their true source (Grassini & Koivisto, 2024)”.

¹¹ <https://openai.com/index/sora>

¹² <https://www.reddit.com/r/brasil/comments/1cml4sj/comment/1311m1h/>

When digital photography appeared, its quality was not comparable to analogue photography. To this day, the texture and dynamic range of chemical images attract a large number of practitioners, even considering their much higher cost and impracticalities. The widespread adoption of numeric imaging, nonetheless, was also responsible for a shift in the habits and the use of photography itself, shifting from memory items “toward communication and identity” (Sarvas & Frohlich, 2011). Another shift is likely to occur, as artificial content starts to become commonplace.

Finally, a last possible cause of society’s rejection of generative content is simply the fact that the Internet is already becoming oversaturated with it. The “Dead Internet Theory”, which began as a conspiracy theory in the early 2020s, is becoming a reality: much of the results returned to Google searches are now AI-generated pages, complete with dull text and “slop” images (Walter, 2024). This could possibly lead to a general scepticism of online content, if not even an utter rejection of the Internet the way we know it today.

5. Scenarios

We would like now to propose some future scenarios to deal with the unfolding of all these challenges presented in the article. These speculations are not mutually exclusive. They are developed as mere fictional thought experiments to extrapolate hypotheses regarding the future of audiovisual creation in response to certain conditions posed with respect to the development of AI. They are hyperboles which, if looked at individually, are highly unlikely to happen, and as such, are not meant to be read as plausible predictions. If anything, we can see that these situations coexist to a small degree as interactions in different territories.

5.1. *TESCREAL Control*

In this scenario, the regulations proposed by AI behemoths like OpenAI were endorsed by governments and codified into law (Kang, 2023). Regulatory desires of tech eschatologists, coinciding with oligopoly interests from big tech, created the perfect lobbying group, pushing populist legislators into creating a restrictive legal framework. In this environment, users are only allowed to access AI tools through the services of companies that, despite having the blessing and authorization of the government, are mostly self-regulated and have no external auditing. All processes are closed under the justification that gatekeepers must keep technology from falling into the wrong hands.

5.2. *Corporate Takeover*

Media conglomerates, allied with artists, manage to promote an even more restricted interpretation of copyright. No material can be used for training without specific permission from owners. There is a commercial war for content, and the models become highly specialized, dependent on the licenses obtained by each company. The US Copyright Office changes its position and allows AI-generated content to be protected. Users must pay for each separate service, making them, in practice, out of reach for the common citizen. Applications become limited to professional use within companies that can afford their costs. Large intellectual property holders profit immensely from the process. Artists obtain a Pyrrhic victory, since the compensation received by each one is insignificant, except for a small elite which is able to negotiate rights individually.

5.3. *Pirate AI*

As the use and storage of models becomes illegal, users band together to exchange files in air-gapped computers, memory sticks, or in encrypted networks. Enforcement agencies begin a war on models, which are accused of enabling terrorism and child pornography. Taking advantage of the situation, rogue countries pass permissive AI laws to attract disreputable businesses. The small Caribbean island of Anguilla, using wealth obtained from licensing the top-level domain *.ai*, positions itself as the free haven for machine learning datasets and open tools.

5.4. *Public AI*

Unwilling to obstruct the technological benefits brought about by artificial intelligence, a consortium of countries is formed to promote publicly trained free models. A pool of collectively managed, publicly available computer hardware becomes the most powerful training system in the world. Legislators reach a compromise regarding copyrights: while no licensing is needed to use proprietary assets, the commercial use of generated materials is not allowed. Even though both artists and media companies are unhappy with the deal, there is an explosion of independent, not-for-profit creativity, enabling a whole new wave of artistic movements. Public science progresses enormously, taking advantage of the power of universally trained large language models and the availability of data in general. This scenario is very similar to the efforts being made by the European Union.

5.5. *Dead Internet*

As the production of generative content becomes ubiquitous, social media companies decide not only to embrace the production of AI “slop” (generic mass-generated AI media) by users, but also to set up their own shop to develop and incorporate artificial characters as users. At first this increases engagement with the networks, but as time passes, this content starts to inevitably appear in new training sets scraped from the web. This has the effect of making new models even more “sloppy”, starting a recursive loop that makes the Internet look increasingly bleak. Finally, a disperse group of committed hackers develop robots that engage with these carcass accounts, hammering the final nail into social networks, and turning the online landscape into a desert devoid of people.

6. Conclusions

The exploration of challenges concerning the input, process, and output of generative audiovisual content highlights the multifaceted and complex nature of these challenges. The legal battles over the use of copyrighted material for training models exemplify the clash between traditional notions of intellectual property and the unprecedented scale of AI-driven content creation. The question of authorship in AI-generated works raises profound philosophical and societal questions about creativity and ownership. As the boundaries between human and machine contributions blur, it becomes important to define clear frameworks for attribution and compensation. Meanwhile, the control over software and data, as well as concerns about the existential risks of AGI, highlight the need for transparency and responsible development practices.

Looking ahead, the scenarios presented offer glimpses into possible futures shaped by different regulatory, economic, and cultural forces. Whether it's a world dominated by corporate interests, a grassroots movement for open access, or a delicate balance between public benefit and private rights, the path forward will require careful navigation and thoughtful consideration of the implications for society as a whole. The scenarios are

clearly exaggerations, but they serve the purpose of preparing audiovisual professionals for the practical obstacles in the path of AI-based creativity. Given the investments made by the major players, the profits expected, and the potential benefits to society, it would be unrealistic to think that the development and adoption of generative AI will be stopped. What might happen is that, as the players involved untangle these problems, new ways of accessing these tools might emerge. These could represent limitations as well as advantages.

Finally, the question of how society will embrace or reject these advances continues to loom over their adoption. AI-generated articles and social media posts are already flooding the Internet, often with little genuine value. As this artificial content becomes more common, the public may start distrusting all user-created content online. Ultimately, users might dismiss any content that doesn't come from established authorities, assuming that it is AI-generated and therefore unreliable.

In the end, the use of AI in audiovisual creation is not just a technological shift but a cultural and ethical reckoning. By engaging in open dialogue, informed policymaking, and collaborative innovation, we can harness the transformative potential of AI while safeguarding the values and principles that define our creative and human endeavours.

References

- Adobe Unveils Future of Creative Cloud with Generative AI as a Creative Co-Pilot in Photoshop.* (2023, May 23). Adobe. <https://news.adobe.com/news/news-details/2023/adobe-unveils-future-of-creative-cloud-with-generative-ai-as-a-creative-co-pilot-in-photoshop>
- Berne Convention for the Protection of Literary and Artistic Works.* (1886). World Intellectual Property Organization. <https://www.wipo.int/treaties/en/ip/berne/index.html>
- Blattmann, A., Dockhorn, T., Kulal, S., Mendelevitch, D., Kilian, M., Lorenz, D., Levi, Y., English, Z., Voleti, V., Letts, A., Jampani, V., & Rombach, R. (2023, November 25). *Stable Video Diffusion: Scaling Latent Video Diffusion Models to Large Datasets.* arXiv.Org. <https://arxiv.org/abs/2311.15127v1>
- Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosselut, A., Brunskill, E., Brynjolfsson, E., Buch, S., Card, D., Castellon, R., Chatterji, N., Chen, A., Creel, K., Davis, J. Q., Demszky, D., ... Liang, P. (2022). *On the Opportunities and Risks of Foundation Models.* arXiv. <http://arxiv.org/abs/2108.07258>
- Broumas, A. (2020). *Intellectual Commons and the Law: A Normative Theory for Commons-Based Peer Production.* University of Westminster Press. <https://doi.org/10.16997/book49>
- Caldas Vianna, B. (2024). *The poetics of autopoiesis: Visual arts, autonomy and artificial intelligence.* Taideyliopiston Kuvataideakatemia. <https://taju.uniarts.fi/handle/10024/8130>
- Calleja Reina, R. (2023). La inteligencia artificial y su derivada en los derechos de propiedad intelectual en la cultura: Retos y amenazas. *Periferica*, 24. <https://doi.org/10.25267/Periferica.2023.i24.06>
- Coeckelbergh, M. (2020). *AI ethics.* MIT Press.
- Copyright, Designs and Patents Act, c. 48, Parliament of the United Kingdom, (3) (1988). <https://www.legislation.gov.uk/ukpga/1988/48/section/9>
- Dean, J. (2022). A Golden Decade of Deep Learning: Computing Systems & Applications. *Daedalus*, 151(2), 58–74. https://doi.org/10.1162/daed_a_01900
- Elbeshbishi, S. (2022, April 24). *Inspiration or infringement? Songwriters clashing in court more often after “Blurred Lines.”* USA TODAY. <https://www.usatoday.com/story/entertainment/music/2022/04/24/copyright-infringement-cases-increase/7367700001/>
- Floridi, L., & Chiriatti, M. (2020). GPT-3: Its Nature, Scope, Limits, and Consequences. *Minds and Machines*, 30(4), 681–694. <https://doi.org/10.1007/s11023-020-09548-1>

- Fuchs, C. (2010). Labor in Informational Capitalism and on the Internet. *The Information Society*, 26(3), 179–196. <https://doi.org/10.1080/01972241003712215>
- Galanter, P. (2003). *What is Generative Art? Complexity Theory as a Context for Art Theory*. 21. <https://www.mat.ucsb.edu/~g.legrady/academic/courses/21s255/r/galanter.pdf>
- Galanter, P. (2016). Generative Art Theory. In C. Paul (Ed.), *A Companion to Digital Art* (1st ed., pp. 146–180). Wiley. <https://doi.org/10.1002/9781118475249.ch5>
- Geburu, T., & Torres, É. P. (2024). The TESCREAL bundle: Eugenics and the promise of utopia through artificial general intelligence. *First Monday*. <https://doi.org/10.5210/fm.v29i4.13636>
- Goodfellow, I. J., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., & Bengio, Y. (2014). Generative Adversarial Networks. *arXiv:1406.2661 [Cs, Stat]*. <http://arxiv.org/abs/1406.2661>
- Grassini, S., & Koivisto, M. (2024). Understanding how personality traits, experiences, and attitudes shape negative bias toward AI-generated artworks. *Scientific Reports*, 14(1), 4113. <https://doi.org/10.1038/s41598-024-54294-4>
- Grynbaum, M. M., & Mac, R. (2023, December 27). The Times Sues OpenAI and Microsoft Over A.I. Use of Copyrighted Work. *The New York Times*. <https://www.nytimes.com/2023/12/27/business/media/new-york-times-open-ai-microsoft-lawsuit.html>
- Heikkilä, M. (2022, September 16). *This artist is dominating AI-generated art. And he's not happy about it*. MIT Technology Review. <https://www.technologyreview.com/2022/09/16/1059598/this-artist-is-dominating-ai-generated-art-and-hes-not-happy-about-it/>
- Hern, A. (2023, March 7). TechScape: Will Meta's massive leak democratise AI – and at what cost? *The Guardian*. <https://www.theguardian.com/technology/2023/mar/07/techscape-meta-leak-llama-chatgpt-ai-crossroads>
- Jiang, H. H., Brown, L., Cheng, J., Khan, M., Gupta, A., Workman, D., Hanna, A., Flowers, J., & Geburu, T. (2023). AI Art and its Impact on Artists. *Proceedings of the 2023 AAAI/ACM Conference on AI, Ethics, and Society*, 363–374. <https://doi.org/10.1145/3600211.3604681>
- Kang, C. (2023, May 16). OpenAI's Sam Altman Urges A.I. Regulation in Senate Hearing. *The New York Times*. <https://www.nytimes.com/2023/05/16/technology/openai-altman-artificial-intelligence-regulation.html>
- Kautz, H. A. (2022). The third AI summer: AAAI Robert S. Engelmore Memorial Lecture. *AI Magazine*, 43(1), 105–125. <https://doi.org/10.1002/aaai.12036>
- Lindberg, V. (2023, February 22). *A mixed decision from the US Copyright Office*. Process Mechanics. <https://www.processmechanics.com/2023/02/22/a-mixed-decision-from-the-us-copyright-office/>
- Liu, J., Li, C., Ren, Y., Chen, F., & Zhao, Z. (2022). DiffSinger: Singing Voice Synthesis via Shallow Diffusion Mechanism. *Proceedings of the AAAI Conference on Artificial Intelligence*, 36(10), Article 10. <https://doi.org/10.1609/aaai.v36i10.21350>
- Metz, C. (2023, November 17). OpenAI's Board Pushes Out Sam Altman, Its High-Profile C.E.O. *The New York Times*. <https://www.nytimes.com/2023/11/17/technology/openai-sam-altman-ousted.html>
- Mittal, G., Engel, J., Hawthorne, C., & Simon, I. (2021). *Symbolic Music Generation with Diffusion Models* (arXiv:2103.16091). arXiv. <http://arxiv.org/abs/2103.16091>
- Mullin, B., & Mickle, T. (2023, December 22). Apple Explores A.I. Deals with News Publishers. *The New York Times*. <https://www.nytimes.com/2023/12/22/technology/apple-ai-news-publishers.html>
- Patterson, D., Gonzalez, J., Le, Q., Liang, C., Munguia, L.-M., Rothchild, D., So, D., Texier, M., & Dean, J. (2021). Carbon Emissions and Large Neural Network Training. *arXiv:2104.10350 [Cs]*. <http://arxiv.org/abs/2104.10350>

- Perrigo, B. (2023, January 18). *Exclusive: The \$2 Per Hour Workers Who Made ChatGPT Safer*. Time. <https://time.com/6247678/openai-chatgpt-kenya-workers/>
- Radford, A., Kim, J. W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., Krueger, G., & Sutskever, I. (2021). *Learning Transferable Visual Models from Natural Language Supervision* (arXiv:2103.00020). arXiv. <https://doi.org/10.48550/arXiv.2103.00020>
- Ren, S. (2023, November 30). *How much water does AI consume? The public deserves to know*. OECD AI Observatory. <https://oecd.ai/en/wonk/how-much-water-does-ai-consume>
- Rombach, R., Blattmann, A., Lorenz, D., Esser, P., & Ommer, B. (2022). *High-Resolution Image Synthesis with Latent Diffusion Models* (arXiv:2112.10752). arXiv. <https://doi.org/10.48550/arXiv.2112.10752>
- Sarvas, R., & Frohlich, D. M. (2011). Digital Photo Adoption. In R. Sarvas & D. M. Frohlich (Eds.), *From Snapshots to Social Media—The Changing Picture of Domestic Photography* (pp. 103–137). Springer. https://doi.org/10.1007/978-0-85729-247-6_6
- Schreiner, M. (2023, July 11). *GPT-4 architecture, datasets, costs and more leaked*. THE DECODER. <https://the-decoder.com/gpt-4-architecture-datasets-costs-and-more-leaked/>
- Shumailov, I., Shumaylov, Z., Zhao, Y., Papernot, N., Anderson, R., & Gal, Y. (2024). AI models collapse when trained on recursively generated data. *Nature*, 631(8022), 755–759. <https://doi.org/10.1038/s41586-024-07566-y>
- Sisson, P. (2024, February 29). A.I. Frenzy Complicates Efforts to Keep Power-Hungry Data Sites Green. *The New York Times*. <https://www.nytimes.com/2024/02/29/business/artificial-intelligence-data-centers-green-power.html>
- Sukhera, J. (2022). Narrative Reviews: Flexible, Rigorous, and Practical. *Journal of Graduate Medical Education*, 14(4), 414–417. <https://doi.org/10.4300/JGME-D-22-00480.1>
- Tong, A., Wang, E., & Coulter, M. (2024, February 22). Exclusive: Reddit in AI content licensing deal with Google. *Reuters*. <https://www.reuters.com/technology/reddit-ai-content-licensing-deal-with-google-sources-say-2024-02-22/>
- USCO. (2021). *U.S. Copyright Office, Compendium of U.S. Copyright Office Practices § 313.2 (3d ed. 2021)*. <https://www.copyright.gov/comp3/>
- Walter, Y. (2024). Artificial influencers and the dead internet theory. *AI & SOCIETY*. <https://doi.org/10.1007/s00146-023-01857-0>
- Warren, S., & Grasser, J. (2024, March 12). Japan Issues Clarification on AI, Copyright, and Machine Learning. *The National Law Review*. <https://natlawreview.com/article/japans-new-draft-guidelines-ai-and-copyright-it-really-ok-train-ai-using-pirated>
- Watson, D. (2019). The Rhetoric and Reality of Anthropomorphism in Artificial Intelligence. *Minds and Machines*, 29(3), 417–440. <https://doi.org/10.1007/s11023-019-09506-6>
- Wolfson, S. (2023, March 23). *The Complex World of Style, Copyright, and Generative AI*. Creative Commons. <https://creativecommons.org/2023/03/23/the-complex-world-of-style-copyright-and-generative-ai/>
- Zeilinger, M. (2021). *Tactical entanglements: AI art, creative agency, and the limits of intellectual property*. meson press.
- Zernay, R. (2017). Casting the First Stone: The Future of Music Copyright Infringement Law after Blurred Lines, Stay with Me, and Uptown Funk. *Chapman Law Review*, 20, 177.