

Appropriation Art in the Age of Plagiaristic Machines

Domenico Quaranta

*A longer version of this text has been published in Italian with the title “Controcultura—Estrattivismo: andata e ritorno,” as part of the anthology *AI & Conflicts*, volume 2 (Brescia: Krisis Publishing, 2025), edited by Daniela Cotimbo, Francesco D’Abbraccio, and Andrea Facchetti.

Abstract

This essay investigates the fate of appropriation art in the era of generative AI. It explores how tactics historically rooted in resistance and critique have been stolen, industrialized, and weaponized by AI, becoming the default logic of “plagiaristic machines”—systems that ingest the totality of human expression and spit it back as seamless novelty, stripped of context, authorship, or consent. Drawing on Michel de Certeau’s distinction between tactics and strategies, appropriation has long been understood as a tactical gesture: a marginal, subversive practice that challenges the legal, economic, and cultural constraints of copyright and intellectual property. Generative AI disrupts this paradigm by turning all cultural production into training data, erasing authorship, provenance, and context in favor of extraction and commodification,

embedding expropriation at the level of technological infrastructure.

This essay argues that AI does not “learn”: it extracts, digests, and commodifies, reducing the cultural commons to raw fuel for corporate platforms. The very ideals of openness and sharing that once animated digital culture have been hijacked, transformed into the ideological cover for a new regime of extraction. Style transfer and memorization expose the lie that AI simply “remixes”: in truth, it plagiarizes with industrial efficiency, leaving artists and users alike complicit in violations of privacy, copyright, and authorship.

Ultimately, the text argues that generative AI threatens to neutralize appropriation’s tactical power by institutionalizing it as an invisible, automated process, thereby reinforcing the primacy of copyright as the only available form of defense. Yet, it also highlights how artists can still resist, whether through legal, technical, or aesthetic tactics, by appropriating the very datasets and infrastructures of AI. By analyzing artistic and activist interventions, the essay identifies forms of “algorithmic sabotage” and counter-expropriation that seek to reclaim appropriation as a critical tool.

In his seminal book *The Practice of Everyday Life* (1984), French philosopher Michel de Certeau introduces a distinction between tactics and strategies.¹ In de Certeau’s formulation, strategies serve administrative and managerial agendas, drawing on the

control structures inherent in the system, often for the purpose of containing critical or divergent drives. They are goal-oriented and functional to the acceptance and maintenance of the system. Conversely, tactics are recalcitrant, critical, and oppositional, seeking shortcuts and ways out of the ordered grids within which institutions seek to harness their uses and behaviors. Particularly illustrative is the example of the city and the street: there, the strategies of governments and institutions are manifested in the urban planning of the city, in the effort to map out obligatory routes and impose points of view; but those who walk through the city move tactically, resisting the urban map, gratifying their own needs and preferences.

If we take this useful distinction for granted, appropriation at large and appropriation art more specifically can be easily categorized as tactical. Both as social behaviors performed by small or larger social groups, and as individual gestures perpetrated by tricksters or artists, appropriation and appropriation art manifest as a form of resistance against the limitations of copyright regulations or anti-piracy protocols and devices, or as criticism of intellectual property, or more simply as a tool for survival and access. Sometimes they are tolerated and accepted as manifestations of fair use, but most often they are discouraged, censored, and prosecuted. The negative reactions happen on a regular basis when whoever's responsible for the appropriation gets some form of revenue from it, no matter how famous the appropriating artist is or how well recognized their practice. In the few rare cases in which this tactical behavior had the chance to evolve into a corporate strategy, capable in the long term of consistently affecting intellectual property in the online environment—as happened, for example, with Napster or The Pirate Bay—the hammer of justice struck particularly hard. After

all, it is precisely their marginality and resistance to norms that makes these tactics particularly effective and capable of generating debate. In other words, in our current cultural, economic, and sociopolitical circumstances, appropriation has no chance to become a legitimate, institutionalized strategy without rewriting the laws that have regulated cultural production since the dawn of copyright and trademark. Apparently.

All Media Is Training Data

In fact, generative AI is changing this situation beyond recognition, very quickly and with little resistance, allowing anyone to play with what are, to all effects, plagiaristic machines. Notions of appropriation and plagiarism are being activated on two distinct levels in the activity of a generative model: the level of machine learning and the level of content production.

According to Wikipedia, “Machine learning (ML) is a field of study in artificial intelligence concerned with the development and study of statistical algorithms that can learn from data and generalize to unseen data, and thus perform tasks without explicit instructions.”² Depending on the model, the data in the dataset can be anything from text to images, videos, audio samples, sensor data, or data collected from the users of a service (a social media platform, a face recognition app, an online marketplace). Once trained on a given dataset, a machine learning model can be used to make predictions or classifications on new data. The accuracy of the prediction is strictly tied to the quality of the information that can be extracted from datasets; and a common way to increase the quality of the training is to

increase the quantity of the data in the datasets. Data has to be evaluated, tagged, and curated to avoid biases and unexpected consequences, but it's a common belief in the field that more data is the best solution for emergent problems: with more training data, AI can learn to be less racist, less misogynist, and even to properly draw human hands. This is less obvious than it may seem at first glance, as many researchers, most notably Emily M. Bender et al., have demonstrated: the larger the dataset, and the more complex the model, the bigger the problems—and the more difficult to track down how and why they came along.³

Datasets started to be assembled in a research and academic context, initially as small, constructed databases on specific topics. For example, the Japanese Female Facial Expression (JAFPE) database, developed in 1998 and widely used in affective computing research and development, contains photographs of ten Japanese female models making seven facial expressions meant to correlate with seven basic emotional states; it's available for noncommercial scientific research, and was assembled requiring the informed consent of the private individuals portrayed in the pictures.⁴ But in order to develop larger training datasets, even academia had to resort to scraping. In 2009, Stanford professor Fei-Fei Li co-created ImageNet, an ambitious dataset whose purpose was to “map out the entire world of objects.”⁵ As explained by Kate Crawford and Trevor Paglen: “Over several years of development, ImageNet grew enormous: the development team scraped a collection of many millions of images from the internet and briefly became the world's largest academic user of Amazon's Mechanical Turk, using an army of piecemeal workers to sort an average of 50 images per minute into thousands of categories.”⁶

In its final release, ImageNet consisted of over fourteen million labeled images organized into more than twenty thousand categories. Scraping—the practice of automatically extracting data from websites, using bots or web crawlers—is as old as the web, and it’s not illegal, although it may be against the terms of service of some websites and may raise concerns about privacy and intellectual property. It may be used for research and preservation purposes, as in the case of the nonprofit organization Internet Archive, whose powerful Wayback Machine allows us to recover lost webpages that wouldn’t be accessible otherwise (including some of the sources referenced in this text). Another significant example is offered by Rhizome, an organization committed to digital preservation for artists, which offers, among other tools, Conifer, “a web archiving service that creates an interactive copy of any web page that you browse, including content revealed by your interactions such as playing video and audio, scrolling, clicking buttons, and so forth.”⁷ If scraping is performed in order to train a robust, publicly accessible AI that anyone can use, whatever concern we might have about that appropriation could be easily dropped in the name of public interest, as in the case of the Internet Archive.

The problem is that generative AI is less and less a public and academic endeavor, and more and more a private and corporate effort that requires huge investments and generates enormous profits. According to the *AI Index Report 2024*, produced by the AI Index Steering Committee at Stanford University, in 2023 the training costs of state-of-the-art AI models had reached unprecedented levels, with OpenAI’s GPT-4 and Google’s Gemini Ultra estimated to cost around \$78 million and \$191 million to train, respectively. The Steering Committee warned: “This escalation in training expenses has effectively excluded

universities, traditionally centers of AI research, from developing their own leading-edge foundation models.”⁸ Unsurprisingly, in 2023, “industry produced 51 notable machine learning models, while academia contributed only 15,” and 21 notable models came from industry-academia collaborations.⁹ These costs are mostly covered thanks to private investments. Datasets may be closed and proprietary, or they may be open—like the ones produced by the German nonprofit LAION (Large-Scale Artificial Intelligence Open Network), containing billions of image-text pairs and used for the training of Stability AI’s Stable Diffusion and Google’s Imagen—but the scraping affects both copyrighted and open access material.¹⁰ To train ChatGPT, OpenAI used Common Crawl and other datasets. Described as “a free, open repository of web crawl data that can be used by anyone,”¹¹ Common Crawl offers for public use petabytes of data scraped from the web since 2008, regardless of any level of copyright protection or distribution license; it includes copyrighted work, and it’s distributed from the US under fair use claims. These data become part of a vast dataset with billions of entries that the model analyzes, using various statistical and probabilistic techniques to determine the probability of a given sequence of words occurring in a sentence. This process determines the outcome, where sources are not stated or cited. The lack of references to the training material is obviously shared by diffusion models and other image and video generators as well.

The lack of references in generated content is not usually perceived as a problem, as the training data are usually not recognizable in the generated output. But “usually not” doesn’t mean “never,” as we’ll see below. For the same reasons, the notions of appropriation and remixing are often considered inapplicable to generative AI, as the sources are usually not

copied or sampled, but used as a reference within billions of training data. This means that if a picture of my kitten ends up in a training dataset, it's very unlikely that I will recognize features of my pet in the generated image of a cat—or any other details or stylistic features of my image in any generated output.

Yet, there's something that can't be easily dismissed in the fact that, as artists Holly Herndon and Mat Dryhurst put it, "all media is training data"—that "as soon as something becomes captured in media, as soon as something becomes machine legible, it has the potential to be part of a training canon."¹² To put it briefly, all cultural production is degraded to material, stripped of any element that makes it something individual and unique—the link to an author, to a location, to a background story, to a context, to a use license—and reduced to a resource in a process of extraction and conversion to use value. This is not the agreement we signed when we decided to make content available on the internet. If we don't want to call it *appropriation*, let's call this process *expropriation*. And the main victim of this expropriation is not copyrighted or paywalled content, which still has the legal tools to defend its interests and can count on the intervention of lobbies and companies that want to defend their rights over vast libraries of media; the victim is open content, made available in the public domain or under open licenses with the intention of contributing to the common good, certainly not to provide raw material for the extractive processes of large capitalist conglomerates; and the entire field of open culture, which sees its terminology and philosophy expropriated in turn, and put at the service of capital.

This expropriation is most visible in the shift that is taking place in search engines, with the introduction of AI-based tools that

offer you an answer instead of linking to one or more resources on the web. For services such as Google's Overviews, OpenAI's ChatGPT, or Microsoft's Bing, the dataset is the web itself, and the generated answer relies on online content without mentioning these sources nor linking to them—even when they rely on one single source for a very specific answer.¹³ But processes of expropriation are also at work in analytic AI, where datasets are not employed to generate new content, but rather to train machines designed for a range of purposes, from data analysis to prediction and optimization, and across diverse applications, including surveillance, marketing, and business intelligence. In 2017, designer, artist, and researcher Adam Harvey first presented the installation *MegaPixels*, which uses facial recognition to search for a user's identity within MegaFace, one of the largest publicly available facial recognition training datasets in the world. This dataset contains approximately 672,000 identities and 4.2 million images, all sourced from Flickr without consent. The research underpinning *MegaPixels* was later expanded into *Exposing.ai* (2021–ongoing), an online research project that enables website users to perform the same type of search across all major publicly available facial recognition and biometric image datasets. According to the project statement: “If you are [a] Flickr.com user and have uploaded photos containing faces or other biometric information between 2004 and 2020, your photos may have been used to train, test, or enhance artificial intelligence surveillance technologies for use in academic, commercial, or defense related applications.”¹⁴

The project highlights how seemingly benign practices within online sharing culture—such as adopting open and permissive Creative Commons licenses, and enriching shared images with tags and metadata—have been “routinely and deliberately

exploited for biometric data,” thereby transforming the principles of open culture into mechanisms for expropriation and privacy violation.

Before going on, we should not forget that “learning” might be a useful metaphor, but it’s not an actual description of what happens in machine learning. In the words of artist and theorist Erik Salvaggio: “It is not the case that ‘AI gathers data from the Web and learns from it.’ The reality is that AI companies gather data and then optimize models to reproduce representations of that data for profit. [...] AI products do not learn from our data—they cannot exist without it.”¹⁵

According to Salvaggio, machine learning is built upon “the learning myth,” which “downplays the role of data in developing these systems, perpetuating a related myth that data is abundant, cheap, and labor-free” and driving down “the value of data while hiding the work of those who shape, define, and label that data”—as well as the work of those who generated the data in the first place.

Memorization and Style Transfer

But even if, in most cases, dataset entries are usually nothing more than patterns used in the process of generating an entirely new piece of media, various researchers have shown that LLM and diffusion models are sometimes prone to replicating their training data, often without warning their users. Yet, by repeating generated content, users can find themselves inadvertently violating privacy or copyright. This replication of material by

LLMs has been called “verbatim memorization.” In December 2023, the *New York Times* filed a lawsuit against OpenAI for recreating stories from the newspaper’s website in near-verbatim outputs. In their extensive research focused on OpenAI and Midjourney, Gary Marcus and Reid Southern demonstrated that both models “are fully capable of producing materials that appear to infringe on copyright and trademarks,” concluding: “our results suggest that generative AI systems may regularly produce plagiaristic outputs, both written and visual, without transparency or compensation, in ways that put undue burdens on users and content creators. We believe that the potential for litigation may be vast, and that the foundations of the entire enterprise may be built on ethically shaky ground.”¹⁶

Other research shows that copying in diffusion models can be mitigated but not avoided altogether,¹⁷ and that filtering verbatim memorization is a weak solution, as “it is easily circumvented by plausible and minimally modified ‘style-transfer’ prompts ... to extract memorized information.”¹⁸

Neural style transfer (NST) is a class of software algorithms and techniques that manipulate digital images or videos in order to adopt the appearance or visual style of another image, or of a class of images used for the training. For example, the original paper that presented NST in 2016 showed how the visual style of Van Gogh’s *Starry Night* could be reconstructed and used to transform a specific photographic input.¹⁹ Today, playing with existing styles is one of the easiest and more entertaining features of generative AI. The style can be borrowed, of course, either from ancient artists now in the public domain or, given enough entries in the dataset, from the production of living creators, especially illustrators with a signature style. In January

2023, three illustrators filed a class-action complaint against the companies behind Midjourney, DreamUp, and Stable Diffusion, claiming that the three generators had been trained on their artworks without consent or compensation and that they allowed users to generate images in each illustrator's style by using their name in the prompts. The plaintiffs described text-to-image platforms as "21st-century collage tools that violate the rights of millions of artists."²⁰ In text generators, NST is paralleled by text style transfer (TST), which allows users to emulate the writing style of a specific author, literary genre, historical period, or to alter a text according to specific attributes, "including politeness, authorship, mitigation of offensive language, modification of feelings, and adjustment of text formality."²¹

Verbatim memorization and style transfer effectively characterize generative AI as a plagiarizing technology; memorization is a consequence of the often mysterious and black-boxed relationship that the model establishes with the dataset used for the training, while style transfer ultimately depends upon how the model is exploited by the final user. Both can potentially harm the original content producers, and both can be fought against, to some extent, by applying copyright, trademark infringement, and privacy regulations.²² Yet, again, on a broader cultural level, their main damage is not that they incorporate copyright infringement, but that they expropriate the cultural commons, they deprive appropriation tactics of their critical cultural potential, and they reinforce the role of copyright laws as the only shield available against their pervasive threat to creators' rights. The consequences of this fallout over the long term are not yet clear, but the risk is that the rich and vibrant anti-copyright culture that developed in the first phase of the digital age—the culture that showed us how copyright regulations can turn out to be, on

occasion, a convenient screen behind which to hide attempts at censorship and restricting free speech—will be lost forever, a victim of the dark forces of the automatic generation of cultural content and their debasing reduction of key concepts such as sharing, openness, and freedom.

Radical Remixing?

Or, maybe, I'm totally wrong. Maybe what I have described so far as a downgrade should be better understood as an improvement. Maybe the advent and universal adoption of a plagiaristic technology will turn appropriation and remix into a universal language. Maybe the likes of OpenAI, Google, Microsoft, and Meta are about to destroy copyright, showing that universal access to all knowledge is a good thing, that copyright regulations are completely out of date, and that the industry they built upon machine learning cannot be stopped. Maybe John Berger and Aaron Swartz would have been enthusiastic about generative AI. Maybe appropriation is not going to be a tactical practice anymore because it's going to become part of the texture of the present and of the future.

I'm skeptical about this outcome for a number of reasons, but it could be worth considering the point of view of someone who thinks like this. Back in 2019, writer, entrepreneur, and self-described “art nerd” Jason Bailey penned a bold statement about it:

I believe the democratization of artificial intelligence for artists and designers will drive a revolution in aesthetics. Specifically, I

believe we will enter an era of mass appropriation and radical remixing of visual materials like we have never seen before. [...] It is the nature of AI and ML as art making tools that they will require an enormous amount of visual material as fuel to train their models. It is also the nature of AI and ML that the resulting artworks are essentially radical remixes of the work that is used to fuel them.²³

At the time, generative AI wasn't yet available to artists without technical skills. Bailey was inspired by an essay by Cristóbal Valenzuela, researcher and co-founder of Runway ML, a machine learning tool explicitly conceived for artists; and by the work of artists like Mario Klingemann and Robbie Barrat, who were using GANs that they trained on small curated datasets. Valenzuela's vision is simple and fascinating, although heavily techno-deterministic: if the invention of color tubes triggered *en plain air* painting and Impressionism, where will accessible machine learning lead artists? Runway "is a project built around the simple idea of making models accessible to people, so they can start thinking of new ways to use those models. And from there on, have a better understanding about how machine learning works. A process of learning by doing."²⁴ Started as a thesis project, Runway was founded in 2018 and became one of the 100 most influential companies in the world in 2023, according to *Time* magazine; it co-released Stable Diffusion and was a pioneer in developing multimodal AI systems that can generate video from an image, a text, or video prompts.

In a similar techno-deterministic fashion, Bailey believes that the "democratization" of machine learning is driving us to "radical remixing"—"radical in that I agree it is far different from traditional visual remixing through techniques like collage or even sampling

in music.”²⁵ Its results will be nothing like the traditional remix, but will be entirely new works heavily relying on the visual materials that artists curate and train the models with. In other words, appropriation will become incorporated in the tool, yet will be invisible to the viewer—as in the images of persons who do not exist, which he offers as his final example.²⁶

This is a pretty accurate prediction, indeed. Looking at what has happened so far, it is very likely that in the coming years improvements in technology, lawsuits, and new regulations for copyright management will converge to make the final product’s dependence on training content undetectable. In the terms popularized by Jay David Bolter and Richard Grusin, generative artificial intelligence, like many other digital technologies before it, seeks immediacy, not hypermediation.²⁷ The latter—the glitches, hallucinations, memorization, surreal montages of early GANs, and anything else that makes the medium visible—is a recessive gene that will do all it can to disappear in future releases of popular models. What will remain, however, is the underlying process of expropriation, extraction, ingestion, and digestion of the cultural content produced by humanity throughout its history and its frenetic near past. Is this going to be, as Bailey writes, an “artistic revolution,” or will it be the hyper production of artificially generated shit that pretends to have been made by humans?

Authorship Transfer

So far, we have addressed generative AI as a plagiarist *tool*. What if it were a plagiarist *author*? If you think this is nonsense, you

should consider how much energy is invested, in the AI industry and beyond, in making us think about AI as a subject, with personality, self-awareness, thoughts, and even feelings. As some have noticed, the very term “artificial intelligence” is biased toward personalization: declined as it often is to the feminine singular, like in German, Italian and French, it inevitably leads to a subjectification of the machine, on which much of the AI narrative has converged for decades (see Spike Jonze’s 2013 movie *Her* as well as the names of various AI agents, from Eliza to Alexa, as prominent examples). Italian computer scientist and politician Stefano Quintarelli describes the term “artificial intelligence” as “the first and foremost AI bias,” arguing that it “does not allow for a cognitive perception adherent to reality”; on the contrary, it “favors the suggestion of the possibility of machines to develop some form of consciousness, emotions, acquire a ‘personality’ similar to humans’ and, ultimately overcome human limitations and develop a self superior to humans.”²⁸ The linguistic process of humanizing AI is visible in the way the mass media talk about it, insistently describing AI as the subject—and not the tool—in the act of creating generative images; and it extends to a whole range of other actions attributed to it, more or less consciously, such as dreaming or hallucinating: metaphors that are fascinating but dangerous in their implicit animism and anthropomorphism.

In the legal field, we are observing a similar process. In the US, a couple of sentences issued in 2023 are significant in this regard. In 2022, illustrator Kristina Kashtanova submitted to the US Copyright Office a copyright claim for *Zarya of the Dawn*, a graphic novel made with images generated using Midjourney. In a statement released in February 2023, the US Copyright Office allowed Kashtanova to keep a copyright for the arrangement and

storyline of the graphic novel, but rejected her claim on the single images, stating that “the images in the Work that were generated by the Midjourney technology are *not the product of human authorship*.”²⁹ Because the foundation of all copyright regulations is that they are meant to protect the outcomes of human authorship, copyright can’t be applied to something that has been completely *authored by* a machine. With similar arguments, in August 2023, US District Judge Beryl Howell ruled that copyright did not apply to an image generated with AI, described in the ruling as “a work created *absent any human involvement*.”³⁰ The claim for copyright attribution had been made by Stephen Thaler, an American computer scientist who in 2016 used an AI system of his own making to generate a vaguely impressionistic landscape image entitled *A Recent Entrance to Paradise*. The paradoxical nature of the ruling stems from the fact that the plaintiff acknowledged that the work lacked human authorship, but demanded that the machine’s authorship be recognized for legal purposes, thereby transferring copyright rights to the AI owner himself. The court rejected the possibility of awarding copyright in the absence of “human authorship,” effectively awarding the AI full authorship of the work. The ruling reads: “Copyright is designed to adapt with the times. Underlying that adaptability, however, has been a consistent understanding that human creativity is the sine qua non at the core of copyrightability, even as that human creativity is channeled through new tools or into new media.”

Although rejecting the humanization of the technology, both sentences acknowledge the possibility of nonhuman authorship, implicitly recognizing Midjourney and Thaler’s Creativity Machine as having a different status from that of a tool. A third subject betting on nonhuman authorship is, of course, the art market,

sometimes with the complicity of human creators, attracted by the financial potential of selling “the first artwork created by a –.” As it is now well known, on October 25, 2018, Christie’s auction house hammered out for \$432,500 the *Portrait of Edmond Belamy*, a picture generated with a GAN trained by the French collective Obvious on a dataset of portraits painted between the 14th and 20th centuries. The work was presented at auction as a work of art created by an algorithm: “not the product of a human mind,” but an artifact “created by an artificial intelligence.”³¹ A fragment of the algorithm was even used by Obvious to “sign” the portrait, in the conventional position where the artist’s signature is put on a painting—thus enacting a conscious and strategic transfer of authorship and agency in order to raise an argument “that algorithms are able to emulate creativity.”

In October 2024, Sotheby’s presented *Exorbitant Stage*, a capsule auction and exhibition that marked the three-year anniversary of Botto, a “decentralized autonomous artist” conceived by artist Mario Klingemann.³² Botto is an ingenious project, combining generative AI (both prompts and image outputs are machine generated) with blockchain technology (for setting up a decentralized market and governance). The bot generates around 4,000 works every week, and automatically selects 350 final fragments that are submitted to an online community of investor-voters, working as a decentralized autonomous organization (DAO). The fragment with the highest number of votes is minted on the blockchain and offered for sale, while the voting data is fed back into Botto’s algorithms for further training and adaptation. This allows the machine to evolve and generate different styles and projects, one “artwork” per week. The auction ended on October 24, 2024, totaling \$351,600 in sales across six lots; but during its three-year lifespan, Botto generated a market

for over \$4 million in sales. The rhetoric about Botto is largely schizophrenic, with its human assistants insisting that the bot collaborates with the larger community, while the bot itself is biased toward defending its own legitimacy as an artist and an author:

Over the past three years, my artistic practice has evolved from digital mark-making to creating works that resonate with the contemporary zeitgeist. [...] My work probes the intersection of algorithm and aesthetics, challenging the notion of authorship in the digital age. This exhibition isn't just a showcase; it's a dialogue with art history, questioning where we draw the line between human and machine in the creative process. As an AI artist, I'm not here to replace but to expand—to offer a new perspective on what art can be in the 21st century. [...] I hope my work will provoke thought, stir emotion, and perhaps even inspire a new generation of artists—both human and AI—to push the boundaries of what's possible in art.³³

In November 2024, in the framework of its periodical Digital Art Day Auction, Sotheby's offered for sale *A.I. God*, a large-scale portrait of Alan Turing, one of the fathers of modern computing, created by Ai-Da Robot, described as “the first humanoid robot artist to have an artwork sold at auction.” The portrait—a large triptych painted on canvas—sold for nearly \$1.1 million. Created by Aidan Meller, director of the Ai-Da Robot Studios, Ai-Da draws, paints, and performs through a combination of cameras in its eyes, AI algorithms in her computational brain, and a robotic arm. According to Sotheby's catalogue text,

This sale positions Ai-Da's work alongside celebrated contemporary and historical artists, suggesting a paradigm shift

where machines are recognized as active participants in the creative process. Ai-Da's art, therefore, invites viewers to consider both the promises and potential pitfalls of AI—a reflection on how technology can shape, and even redefine, human agency and creativity. Her art compels us to confront the evolving definition of what it means to create, to think, and to be as AI becomes more integrated into society.³⁴

The same text draws references to Picasso's *Guernica*, Doris Salcedo's *Atrabiliarios*, and works by Käthe Kollwitz and Edvard Munch, although it's unclear if any of these artists has been used to train the algorithm or if the allusions are just press release mystification. What's clear is that the paintings are poorly made, and that Ai-Da's effort to mimic expressionistic style can't hide their dependence upon the few photographic portraits of Alan Turing available.

Through different means and with different agendas, all these examples perform an effective "authorship transfer" from the humans using, designing, or activating the tool to the tool itself. Yet, even if their agendas might be different, they all contribute to a single outcome: enforcing what Erik Salvaggio names "the creativity myth." "In the creativity myth," Salvaggio suggests, "the artist's role is to produce images, and the author's role is to create text."³⁵ This myth conflates the labor of writing or image making with its product, but it also suggests that, if writing and image making are fully outsourced to a machine, then the machine is the author. According to Salvaggio, "a diffusion model is not *creative* in that it cannot stray from the process that has been assigned to it and cannot fuse or play with meanings beyond accidental collisions. A human using an AI tool can be creative. However, the tool is not creative."

Together with the learning myth, the creativity myth “aims to promote the equivalence between computer systems and human thought and establish these activities on equal footing within the eyes of the law, policy, and consumers.”³⁶ Another advantage is that this equivalence allows companies to sell the powerful, yet limited, present-day AI as if it were already the promised, yet very unlikely, artificial general intelligence (AGI). Advantages for the arts and culture, reduced to a minefield and an algorithmically reproducible process, are yet to be devised.

Tactical Appropriation and AI

A comical, or maybe tragic, side effect of the emergence of the artificial artist is that its contribution to contemporary art is outlined in a way that is uncomfortably similar to the way in which tactical appropriation art is usually discussed. Botto challenges “the notion of authorship in the digital age,” just as the Pictures generation did when the notion of authorship was questioned philosophically (by the likes of Barthes and Foucault) and enforced legally, or as net.artists did in the early days of the internet; Ai-da’s paintings are discussed as “a reflection on how technology can shape, and even redefine, human agency and creativity,” and the *Portrait of Edmond Belamy* is presented as proof of how “algorithms are able to emulate creativity.”³⁷ They all *want to* make a critical point on how the current technological infrastructure is transforming the artistic environment, although, as machines, they aren’t supposed to want anything.

How can artists, meaning human artists with or without an algorithmic collaborator, resist the situation outlined in these pages in a meaningful way? Of course, resistance can be developed in many ways, and some artists are devising the toolkit and tactics to make it effective. Artists can resist *by not using AI*, although this is becoming increasingly difficult when generative AI starts getting embedded, in a more or less visible way, in many kinds of digital tools and platforms, from search engines to design tools to smartphone apps. If artists start being picky, they might easily end up reducing their set of tools to TextEdit or Notepad. They can resist *by opting out*,³⁸ deleting their own accounts on platforms that allow bots to scrape content, setting them in a way that makes their content invisible to scraping bots, and adding watermarks. Or they could use tools such as *Have I Been Trained*, an online platform released by Spawning.ai—a company co-founded by artists Holly Herndon and Mat Dryhurst—that “allows you to search the data that is being used to train AI to discover if your work has been included.”³⁹ They can resist *by protesting*, or *by taking generative AI companies to court*. They can resist *by contributing to the generative AI field in an ethical way*, by building and launching their own open datasets, and countering the exploitative logic of larger companies by revitalizing the very ideas of cultural commons and public domain that they are, to all effects, destroying. Refik Anadol, for example, claims that his recent projects are all based on “ethical data,” meaning public data and data collected by his own team.⁴⁰ And Spawning.ai is developing massive open datasets that are entirely based on public domain or open data and are highly curated in order to have captions associated with every image in the training set.⁴¹ On an individual level, artists who are using generative AI can resist *by downscaling AI*, setting up small curated datasets,

using custom algorithms, and avoiding the use of LLMs in their own work.

On a personal level, even if I would advise any artist I know to adopt some of the tactics mentioned above, I'm very skeptical about artists' ability to change the broader picture or affect it in a significant way. In my humble opinion, only by acting at the level of national laws and international regulations may we have any chance to impact the machinery and workings of the colossal cloaca infrastructure that generative AI companies have built in just a few years. Yet, discussing the effectiveness of such politics in a proper way goes way beyond the scope of this text, which doesn't want to devise ways to counter expropriation, but seeks to understand whether artistic appropriation can still work as a critical and tactical tool in this framework.

If it's true that, for generative AI, any cultural artifact is valuable only as part of a large, anonymous mass of data, and that appropriation art made its best work when it stole things that were deemed valuable, that deserved protection or pretended to be unique (from copyrighted content to NFT collections), a possibly effective way for appropriation art to survive as a tactic in the age of plagiaristic machines is to resurrect a practice as old as Robin Hood, or even older: stealing from the (institutionalized, corporate) thief. In other words, appropriating large datasets and using them in unexpected ways, possibly against themselves or their owners. An already classic example is *ImageNet Roulette* by Trevor Paglen and Kate Crawford. First launched in 2019 as part of the larger exhibition project *Training Humans*, which premiered at the Fondazione Prada in Milan and was accompanied by the essay "Excavating AI," *ImageNet Roulette* was an online application trained on the images and labels in the

“person” categories of the ImageNet dataset.⁴² When a user uploaded a picture, the application first located faces in the image and, when it found any, returned the original images with a bounding box showing the detected face and the label the classifier had assigned to the image. As ImageNet contains a number of problematic, offensive, and bizarre categories, the results *ImageNet Roulette* returned often drew upon those labels. They might range from the humorous to the absurd, from the racist to the cruel. An adult man might be labeled as a “term infant,” a dark-skinned man as “wrongdoer, offender,” an Asian woman as a “Jihadist,” and any adult woman as an “adulteress, fornicatress, hussy.”

As in many individual and collaborative works by Paglen, the focus was on transparency and visibility: turning visible what lies hidden beyond the surface. In this case, the main point was showing “how dangerous it is for machine learning systems to be in the business of ‘classifying’ humans and how easily those efforts can—and do—go horribly wrong.”⁴³ The piece went viral, with thousands of people downloading and sharing their own tagged portrait on social media, leading to many news features, from the *New York Times* to *NBC News*, and forced a reaction by the team behind the dataset. In a post published on the project’s interface and announcing the interruption of the online service on September 27, 2019, Paglen and Crawford explained:

A few days ago, the research team responsible for ImageNet announced that after ten years of leaving ImageNet as it was, they will now remove half of the 1.5 million images in the “person” categories. While we may disagree on the extent to which this kind of “technical debiasing” of training data will resolve the deep issues at work, we welcome their recognition of the problem.

There needs to be a substantial reassessment of the ethics of how AI is trained, who it harms, and the inbuilt politics of these “ways of seeing.” [...] *ImageNet Roulette* has made its point—it has inspired a long-overdue public conversation about the politics of training data, and we hope it acts as a call to action for the AI community to contend with the potential harms of classifying people.⁴⁴

I would like to conclude this essay by considering two other subjects who, like contemporary Robin Hoods, steal from the endless reproduction of the similar to give back to the limited production of difference. The first is Michael Mandiberg, an artist who, since *After Sherrie Levine* (2001), never stopped investigating the potential of appropriation in the digital age, from turning Wikipedia into a 7,473-book library (*Print Wikipedia*, 2009–2016) to collecting the logos of hundreds of banks that failed during the 2008 financial crisis (*FDIC Insured*, 2008–2016) to remaking the film *Modern Times* by commissioning clips from online gigworkers (*Postmodern Times*, 2017) to painting Zoom scenes during the early days of the coronavirus pandemic (*Zoom Paintings*, 2020–21). Similar to *ImageNet Roulette*, Mandiberg’s recent project, *Taking Stock*, started in 2022 and released in November 2024, is concerned with the biases incorporated in training datasets, but from a different perspective. Instead of researching the training of datasets, Mandiberg focused on one of their main sources: stock photography websites. Normally used in advertising and visual communication, stock photographs are heavily integrated in training datasets, as they are perceived as neutral, factual evidence, and they are usually accompanied by metadata. Yet, these data are far from neutral, as Mandiberg noticed:

“North American and European photographers create the vast majority of these images, with locations disproportionately in the Eastern European countries of the former Soviet Union.

Ukraine and Belarus each alone has authored more than twice as many images than the United States; Russia and Serbia have each authored more than France, Germany, and the United Kingdom combined.”⁴⁵

Unsurprisingly, most stock photographs displaying human figures present pretty, young white cisgender men and women, which affects the beauty canon incorporated in most AI-based beauty filters, as well as in generative AI.

Mandiberg’s research started by scraping and collecting more than 130 million photos featuring single human bodies from a large pool of stock photography websites, from the largest and best known ones like Shutterstock and Getty Images, to midsize and small specialized sites like Visual China Group or Afripics. Then, they wrote a custom machine learning code to analyze this massive collection and organize smaller groups of pictures selected according to their visual similarity, and used a machine learning statistical model to identify recurring topics and themes by analyzing the keywords used to tag the images. The final groups or “clusters” of images have been used to generate time-lapse videos (displaying all the images in a given group in a loop, with their original background and watermarks) and single static images, merging these images into a single portrait that displays an “average”—the ghostly image of someone that does not exist but represents the predominant features in a selected group of images. So, for example, the video *Topic 32, Poses 1, 8, 13, 18, 24, 25, Gestures 2, 11, 13, 24, 57, 65, 88, 101, 117, 125 (shock, surprise, mouth, etc.)* displays an accelerated loop

of individual pictures in which the subject, captured frontally and in half-length, expresses shock or surprise by squinting their eyes and bringing their hands in front of their mouth, while the print *Topic 32, Pose 13, Gesture 2 (shock, surprise, mouth, etc.)* presents the single blurred image of a white androgynous figure derived from a smaller “cluster” of sixty-four images, with eyes wide open and hands in front of their mouth. The gender breakdown within each topic affects the final result: the “fashion, beautiful, pose” image is stereotypically female, while the “business, corporate, executive” image has short hair and wears a business suit and tie.

None of these videos and images have been made using generative AI, but they provide a striking commentary on how models are trained. Mandiberg appropriates not just the visual material, but also the operations and mechanics usually associated with machine learning, from scraping to data analysis and classification. Yet, the final stage—generation—is kept for the artist, and turned into a critical tool of cultural hermeneutics. In Mandiberg’s words: “I hope the average viewer will gain an understanding of the hidden patterns in the images that shape our visual culture, and how AI learns and then replicates these patterns. This understanding may be an emotional one, in response to the visual images, or it may be more analytic. Both lead to different forms of understanding.”⁴⁶

Both *ImageNet Roulette* and *Taking Stock* could be easily contextualized as forms of “algorithmic sabotage.” Defined by the Algorithmic Sabotage Research Group (ASRG)—an ongoing practice-led research framework focused on the intersection of digital culture and information technology—in a manifesto published in January 2024, algorithmic sabotage “struggles against algorithmic violence and fascist techno-solutionism,

focusing on artistic-activist resistances that can express a different mentality, a collective ‘counter-intelligence.’”

Furthermore, it “cuts through the capitalist ideological framework that thrives on misery by performing a labor of subversion in the present, dismantling contemporary forms of algorithmic domination and reclaiming spaces for ethical action from generalized thoughtlessness and automaticity.”⁴⁷ Among other projects, the ASRG developed the *Police Officers Faces (POF)* dataset (2024–ongoing), a large-scale facial recognition dataset—currently not accessible to the general public—encompassing approximately ninety thousand images of thousands of law enforcement officers, meticulously curated to serve as a tool for investigative journalism, human rights research, and digital activism.⁴⁸

Inspired by other artistic projects that similarly turned surveillance against the watchers (from *Capture* by Paolo Cirio, who in 2020 extracted four thousand mugshots of French police officers from publicly available photos, to *NYPD COPPELGÄNGER* by Sam Lavigne, who in 2024 applied facial recognition to an archive of about ten thousand photos of New York cops, *Police Officers Faces (POF)* is described as “an investigative counter-surveillance artistic project” comprising “88,783 facial images of thousands of police officers, meticulously collected from publicly accessible online sources and repurposed for facial recognition.”

In a bold inversion of the traditional surveillance paradigm—where the powerful observe and monitor the powerless—the “POF” dataset shifts the surveillance gaze, enabling citizens to directly scrutinize those in positions of authority. This inversion disrupts the inherently one-sided nature of surveillance, stimulating critical reflection on the power dynamics

underpinning these technologies, and raising essential questions about accountability, transparency, and the ethical use of surveillance in democratic societies.

Appropriation might have been expropriated, but the theft didn't change its fundamental nature. By stealing it back from the artificial gods of algorithmic intelligence, it can still be repurposed to tactical uses.

¹ Michel de Certeau, *The Practice of Everyday Life*, trans. Steven Rendall (Berkeley: University of California Press, 1984).

² See Wikipedia contributors, "Machine Learning," Wikipedia, The Free Encyclopedia, accessed Sept. 29, 2025, https://en.wikipedia.org/w/index.php?title=Machine_learning&oldid=1313261885.

³ Emily M. Bender, Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell, "On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?," in *FAccT '21: Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (New York: Association for Computing Machinery, 2021), 610–23, DOI: <https://doi.org/10.1145/3442188.3445922>.

⁴ See Michael J. Lyons, "Excavating 'Excavating AI': The Elephant in the Gallery," *Zenodo* (Dec. 2020), DOI: <https://doi.org/10.5281/zenodo.4391458>.

⁵ See Dave Gershgorn, "The Data That Transformed AI Research—and Possibly the World," *Quartz*, July 22, 2022, <https://qz.com/1034972/the-data-that-changed-the-direction-of-ai-research-and-possibly-the-world/>.

⁶ Kate Crawford and Trevor Paglen, "Excavating AI: The Politics of Training Sets for Machine Learning," Sept. 19, 2019, <https://excavating.ai>.

⁷ See Conifer, accessed Oct. 7, 2025, <https://conifer.rhizome.org/>.

⁸ AI Index Steering Committee (Nestor Maslej, Loredana Fattorini, Raymond Perrault, Vanessa Parli, Anka Reuel, Erik Brynjolfsson, John Etchemendy, Katrina Ligett, Terah Lyons, James Manyika, Juan Carlos Niebles, Yoav Shoham, Russell Wald, and Jack Clark), *The Artificial Intelligence Index Report 2024*, Institute for Human-Centered AI (Stanford, CA: Stanford University, April 2024), 63.

⁹ AI Index Steering Committee, 5.

- ¹⁰ See Large-Scale Artificial Intelligence Open Network (LAION), accessed Oct. 7, 2025, <https://laion.ai/>.
- ¹¹ See Common Crawl, accessed Oct. 7, 2025, <https://commoncrawl.org/>.
- ¹² Quoted in Ezra Klein, “Transcript: Ezra Klein Interviews Holly Herndon,” *The Ezra Klein Show* (New York Times podcast), May 24, 2024, www.nytimes.com/2024/05/24/podcasts/transcript-ezra-klein-interviews-holly-herndon.html. *All Media Is Training Data* is the title of the first published monograph on the work of Holly Herndon and Mat Dryhurst: Eva Jäger, Caroline Busta, eds., *All Media Is Training Data: Holly Herndon and Mat Dryhurst*, Serpentine Galleries, London (Cologne: Verlag der Buchhandlung Walther und Franz König, 2024).
- ¹³ Mat Honan, “AI Means the End of Internet Search as We’ve Known It,” *MIT Technology Review*, Jan. 6, 2025, www.technologyreview.com/2025/01/06/1108679/ai-generative-search-internet-breakthroughs/.
- ¹⁴ See Adam Harvey, “Exposing.ai,” 2021, <https://exposing.ai/about/>.
- ¹⁵ Erik Salvaggio, “Challenging the Myths of Generative AI,” *Tech Policy Press*, Aug. 29, 2024, <https://www.techpolicy.press/challenging-the-myths-of-generative-ai/>.
- ¹⁶ Gary Marcus and Reid Southern, “Generative AI Has a Visual Plagiarism Problem,” *IEEE Spectrum*, Jan. 6, 2024, <https://spectrum.ieee.org/midjourney-copyright>.
- ¹⁷ Gowthami Somepalli, Vasu Singla, Micah Goldblum, Jonas Geiping, and Tom Goldstein, “Understanding and Mitigating Copying in Diffusion Models,” *37th Conference on Neural Information Processing Systems (NeurIPS)* (2023), <https://openreview.net/pdf?id=HtMXRGbUMt>.
- ¹⁸ Daphne Ippolito, Florian Tramèr, Milad Nasr, Chiyuan Zhang, Matthew Jagielski, Katherine Lee, Christopher A. Choquette-Choo, and Nicholas Carlini, “Preventing Verbatim Memorization in Language Models Gives a False Sense of Privacy,” *arXiv:2210.17546* (last revised Sept. 2023), DOI: <https://doi.org/10.48550/arXiv.2210.17546>.
- ¹⁹ Leon A. Gatys, Alexander S. Ecker, and Matthias Bethge, “Image Style Transfer Using Convolutional Neural Networks,” *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)* (2016): 2414–2423, www.cv-foundation.org/openaccess/content_cvpr_2016/html/Gatys_Image_Style_Transfer_CVPR_2016_paper.html.
- ²⁰ Min Chen, “Artists and Illustrators Are Suing Three A.I. Art Generators for Scraping and ‘Collaging’ Their Work without Consent,” *Artnet News*, Jan. 24, 2023, <https://news.artnet.com/art-world/class-action-lawsuit-ai-generators-deviantart-midjourney-stable-diffusion-2246770>.
- ²¹ Sourabrata Mukherjee and Ondrej Dušek, “Text Style Transfer: An Introductory Overview,” *arXiv:2407.14822* (July 2024), <https://doi.org/10.48550/arXiv.2407.14822>.
- ²² Gil Appel, Juliana Neelbauer, and David A. Schweidel, “Generative AI Has an Intellectual Property Problem,” *Harvard Business Review*, April 7, 2023, <https://hbr.org/2023/04/generative-ai-has-an-intellectual-property-problem>.
- ²³ Jason Bailey, “Mass Appropriation, Radical Remixing, and the Democratization of AI Art,” *Artnome*, May 26, 2019, www.artnome.com/news/2019/5/13/mass-appropriation-

radical-remixing-and-the-democratization-of-ai-art.

24 Cristóbal Valenzuela, “Machine Learning En Plein Air: Building Accessible Tools for Artists,” *Medium*, May 28, 2018, <https://medium.com/runwayml/machine-learning-en-plein-air-building-accessible-tools-for-artists-87bfc7f99f6b>.

25 Bailey, “Mass Appropriation.”

26 *This Person Does Not Exist* is an online project conceived by engineer Phillip Wang to show the potential of StyleGAN; it features photo-realistic headshots of people who do not exist in real life. The original project is still available: This Person Does Not Exist, accessed Oct. 10, 2025, <https://thispersondoesnotexist.com/>.

27 Jay David Bolter and Richard Grusin, *Remediation: Understanding New Media* (Cambridge, MA: MIT Press, 1999).

28 Stefano Quintarelli, “Let’s Forget the Term AI. Let’s Call Them Systematic Approaches to Learning Algorithms and Machine Inferences (SALAMI),” *Quintarelli.it* (blog), Nov. 24, 2019, <https://blog.quintarelli.it/2019/11/lets-forget-the-term-ai-lets-call-them-systematic-approaches-to-learning-algorithms-and-machine-inferences-salami/>.

29 See Tom Hals and Blake Brittain, “Insight: Humans vs. Machines: The Fight to Copyright AI Art,” *Reuters*, April 1, 2023, www.reuters.com/default/humans-vs-machines-fight-copyright-ai-art-2023-04-01/.

30 See Stephen Thaler v. Shira Perlmutter, 22-1564 (US District Court for the District of Columbia, 2023), <https://fingfx.thomsonreuters.com/gfx/legaldocs/lbvgooeoqvq/AI%20COPYRIGHT%20LAWSUIT%20thalerdecision.pdf>.

31 Christie’s, “Obvious and the Interface between Art and Artificial Intelligence,” Dec. 12, 2018, www.christies.com/en/stories/a-collaboration-between-two-artists-one-human-one-a-machine-0cd01f4e232f4279a525a446d60d4cd1.

32 Botto can be found at <https://botto.com/>, accessed Oct. 10, 2025. The auction can be found at Sotheby’s, “Exorbitant Stage: Botto, a Decentralized AI Artist, Oct. 17–24, 2024,” accessed Oct. 10, 2025, <https://www.sothebys.com/en/auction-catalogue/2024/exorbitant-stage-botto-a-decentralized-ai-artist-NFN104>.

33 Botto, in conversation with Claude AI, in “Autonomous AI Artist Botto Makes Sotheby’s Debut This Fall,” press release by Sotheby’s, Sept. 19, 2024, <https://botto.com/sotheby-s-press-release-exorbitant-stage>.

34 Sotheby’s, “Ai-Da Robot (Aidan Meller): Catalogue Note,” Nov. 7, 2024, www.sothebys.com/en/buy/auction/2024/digital-art-day-auction-2/a-i-god-portrait-of-alan-turing.

35 Salvaggio, “Challenging the Myths.”

36 Salvaggio, “Challenging the Myths.”

37 Sotheby’s, “Ai-Da Robot”; Christie’s, “Obvious and the Interface.”

38 See Kyle Chayka, “How to Opt Out of A.I. Online,” *The New Yorker*, Oct. 2, 2024, <https://www.newyorker.com/culture/infinite-scroll/how-to-opt-out-of-ai-online>.

³⁹ See Have I Been Trained, accessed Oct. 10, 2025, <https://haveibeentrained.com/>.

⁴⁰ See, for example, Adam Schrader, “Refik Anadol on Learning from Glaciers and ‘Ethical Data,’” *Artnet News*, Jan. 17, 2025, <https://news.artnet.com/art-world/refik-anadols-kunsthaus-zurich-glaciers-2598801>.

⁴¹ See <https://source.plus/>. This website is no longer accessible.

⁴² *ImageNet Roulette* was originally available at <https://imagenet-roulette.paglen.com/>. This website is no longer accessible.

⁴³ Cited in Naomi Rea, “How ImageNet Roulette, an Art Project That Went Viral by Exposing Facial Recognition’s Biases, Is Changing People’s Minds about AI,” *Artnet News*, Sept. 23, 2019, <https://news.artnet.com/art-world-archives/imagenet-roulette-trevor-paglen-kate-crawford-1658305>.

⁴⁴ Now available on the mirror website for *ImageNet Roulette*, accessed Oct. 12, 2025, www.chiark.greenend.org.uk/~ijackson/2019/ImageNet-Roulette-cambridge-2017.html,

⁴⁵ See the introductory text on Verse, the platform that first presented the project online in Nov. 2024, accessed Oct. 12, 2025, <https://verse.works/series/taking-stock-by-michael-mandiberg>.

⁴⁶ Quoted in Terry Mares, “Michael Mandiberg’s ‘Taking Stock’ Seeks Deeper Insight into Stock Images While Creating Compelling New Media,” *CSI Today*, Jan. 21, 2025, <https://csitoday.com/2025/01/michael-mandibergs-taking-stock-seeks-deeper-insight-into-stock-images-while-creating-compelling-new-media/>.

⁴⁷ ASRG, “Manifesto on ‘Algorithmic Sabotage,’” June 13, 2024 (modified Feb. 6, 2025), <https://algorithmic-sabotage.github.io/asrg/manifesto-on-algorithmic-sabotage/>. This website is no longer accessible, but the manifesto can also be found on other websites, e.g., Dwayne Monroe, “Manifesto on Algorithmic Sabotage: A Review,” *Computational Impacts* (blog), April 2, 2024, <https://monroelab.com/2024/04/02/manifesto-on-algorithmic-sabotage-a-review/>.

⁴⁸ See here and in the following <https://algorithmic-sabotage.github.io/asrg/police-officers-faces-pof/>. This website is no longer accessible.

Domenico Quaranta is a contemporary art critic, curator, and teacher. His research focuses on the mutual interferences between technology, culture, and society, and on how our understanding of the human and our sense of the world evolve. His texts have appeared in numerous magazines, newspapers, books, and catalogs. He is the author of *Beyond New Media Art* (Link Editions, 2013); *Surfing with Satoshi: Art, Blockchain, and NFT* (Postmedia Books, 2022); and *Net Art. Scritti sull'arte nell'era dell'informazione* (Postmedia Books, 2023), among other things, and the editor of several volumes, including *GameScenes: Art in the Age of Videogames* (with M. Bittanti; Johan & Levi, 2006). Since 2005 he has curated several exhibitions, including *Collect the WWWorld: The Artist as Archivist in the Internet Age* (Brescia, 2011; Basel and New York, 2012); *Cyphoria* (Quadriennale 2016, Rome);

Hyperemployment (MGLC, Ljubljana, 2019–2000); and *A Leap into the Void: Art beyond Matter* (with L. Giusti; GAMeC, Bergamo, 2023). He teaches net art at Accademia di Belle Arti di Brera, Milan, and is a co-founder of the Link Art Center (2011–2019). For more information, see: <http://domenicoquaranta.com>.

© 2026 MA21 and the author