

# The Future of Using Public Images for AI Research

Originally published at <https://arquivo.pt/wayback/20240203194728oe/https://blog.metaphysic.ai/the-future-of-using-public-images-for-ai-research/>

November 8, 2022



Besides their capacity to [rip off](#) the style of popular real world artists, the new breed of [latent diffusion](#)-based image synthesis systems promises a revolutionary ease-of-use not only for valid creative purposes, such as concept art [development](#) and [stock image generation](#), but also for creating [controversial deepfakes](#), and potentially objectionable imagery, including [child pornography](#).

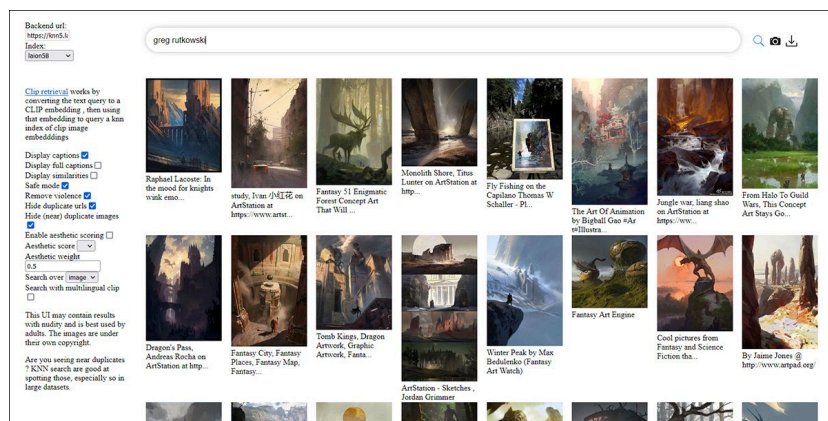
In the ordinary run of things, these would be signals of coming restrictions. Anyone who remembers [the late 1990s](#) will know how the sudden unleashing of unprecedented and unrestricted technologies leads first to proliferation, then public inquiries, and then, eventually, some kind of governance and regulation ([not always successful](#)), forcing these 'edge activities' into the seclusion of private communities or [anonymized platforms](#).



Examples of casual Stable Diffusion submissions from [r/stablediffusion](#).

Therefore it would seem a foregone conclusion that the current 'freedom to fake' or produce any kind of image that you can imagine, as provided by the mercurial and [instant rise](#) of Stable Diffusion, will eventually be curtailed by laws that prevent the exploitation of the web-scraped public images that give image synthesis systems their extraordinary transformative powers.

Around the world, this process is [still developing](#) for the [autoencoder deepfakes](#) that power both viral 'fake celebrity' videos, and deepfake porn.

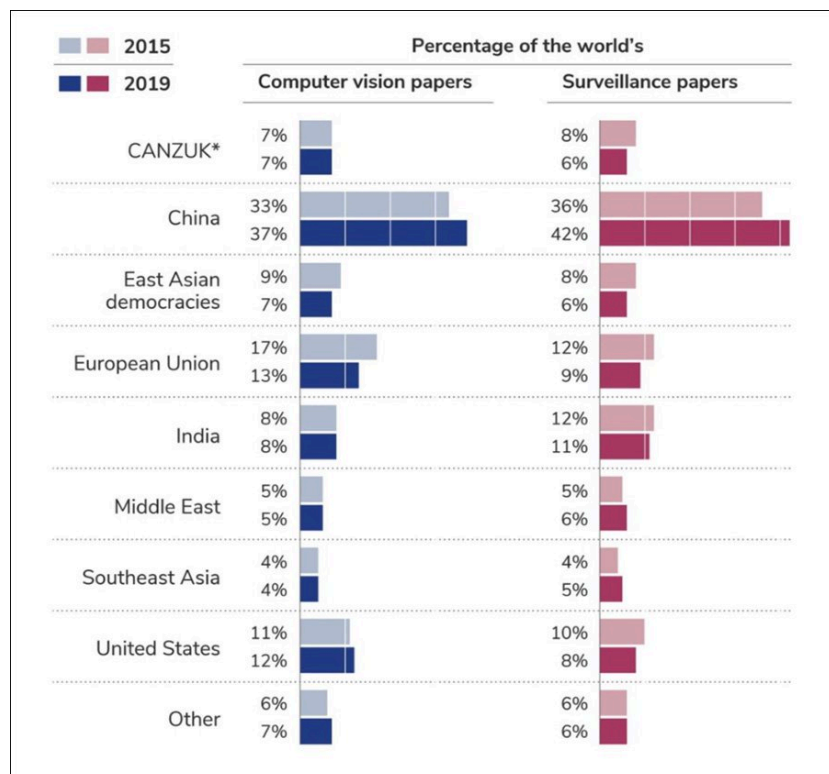


How the Greg Rutkowski craze began – with the inclusion of the artist's work in the LAION- aesthetics database that subsequently trained the models for Stable Diffusion. Source: <https://rom1504.github.io/clip-retrieval/>

## A Special Case

However, in the case of 'locking down' this apparently unlimited font of free, web-found image data so that it can't contribute to the output of image synthesis architectures, there are unusual opposing forces that seem likely to factor into the problem.

Put plainly, western governments are concerned about China's significant and growing advantage in many areas of AI – and, according to a [2022 report](#) from the Center for Security and Emerging Technology (CSET), Beijing has a particularly [pronounced lead](#) in computer vision research, and in the sub-field of surveillance technologies (the majority of which are related to computer vision) – areas where it [lagged](#) the west little more than two years ago.



china-lead

The CSET report states:

*'Globally, research output in computer vision and visual surveillance has grown over time, but China's publication rate rose especially rapidly between 2015 and 2019. As a result, China grew as a contributor to visual surveillance research, from 36 percent of global surveillance research in 2015 to 42 percent in 2019. By contrast, other regions held steady or lost research share.'*

If anything, the report notes, the authors are underestimating the potential extent of China's dominance and capacity to grow in this space. However, they can only gauge metrics from what China chooses to publish.

The only truly affordable, [authentic](#) and sustainable hyperscale source for data-hungry computer vision systems remains the ever-increasing flow of [2.3 billion images](#) uploaded daily to the internet, most of which are liberally accessible to web-scraping systems, and which collectively provide updated and relevant classes and semantic terms, in contrast to older 'anchor' datasets, which were gathered in a different cultural climate and under greater logistical restrictions.

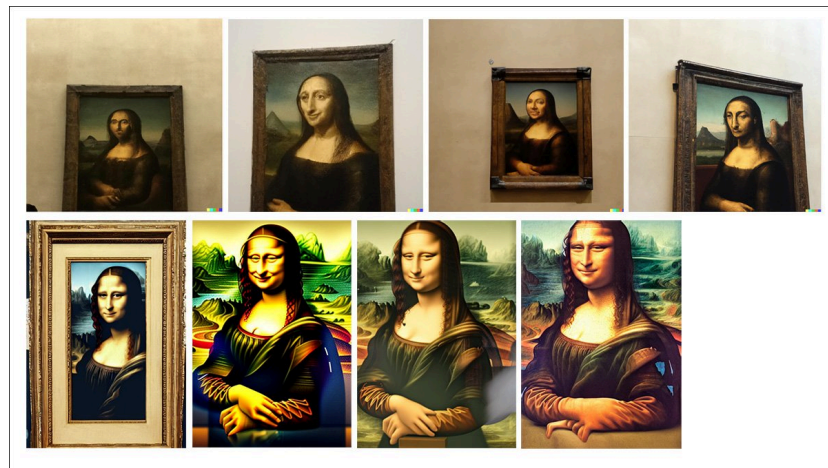
The question coming into focus now is whether or not growing concern about the ongoing exploitation of these images in generative systems can hope to lead to greater regulation in the west, when such laws may further widen the computer vision technology gap with Asia.

## Synthesis or Plagiarism?

"I think that two things will happen in the next 12-18 months," says† Bradford Newman, a litigation partner for global law firm Baker McKenzie, and a [vocal proponent](#) of AI data regulation. "One is that a lawsuit, or multiple lawsuits will be filed which will take years to resolve. The other is that you're going to start seeing major lobbying from AI influencers, as well as think-tanks and other ivory tower players, and private sector players. They're all trying to monetize AI, and they'll be looking for a senator or congresswoman to adopt their cause."

Newman believes that any journey towards additional US legislation is going to be slow and beset with complications – many of them logistical. In the short term, he believes that various systems of rated compensation (such as the one [recently announced](#) by stock image company Shutterstock) for those whose work has contributed to generative systems will be the earliest, though not necessarily the ultimate solution.

"The problem," he says, "is proving causality. When a generative algorithm receives a text-prompt about *sneakers*, and produces an image to which your own images in the dataset were a major contribution, it's hard to prove that relationship. If a generative system won't even reproduce the Mona Lisa exactly, it isn't going to precisely reproduce your work either."



*'The Mona Lisa': The output of DALL-E 2 (above), and various checkpoints and settings in Stable Diffusion (below), none exactly reproducing the art classic. Though this well-known work could probably be identified through facial recognition or other interstitial systems inside a text-to-image architecture, not many images in a hyperscale database are this engrained in the culture, and most contributing photos are 'generalized' into obscurity during model training, making it difficult to know the extent to which any one image may have influenced the output.*

"So it's hard to establish the point at which you ought to get paid – unless we enact a law requiring some kind of generic royalty payment for having made any contribution at all to the dataset, similar to what the music industry received twenty years ago as a tax†† on blank DVDs."

Under such an egalitarian scheme, a text-to-image luminary like painter Greg Rutkowski, whose name is [possibly the most-used term](#) in Stable Diffusion prompts, would receive the same compensation as an obscure 1920s photographer who's rarely ever used in a prompt.

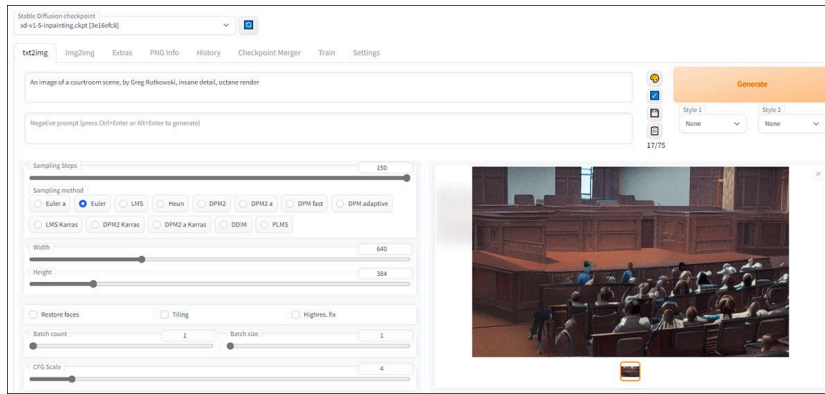
Though this abandonment of merit-based reward is arguably a tad socialist even for European legal principles, much less those of the United States, it may be the only realistic solution in the near-term, even if it would in any case require the disclosure of the use of an image synthesis system in a published work in order to invoke a charge, or else the existence of [detection systems](#) capable of identifying fully synthesized images; or even detecting when a generative system has ['reproduced' an original work](#). Successful detection of this kind is likely to get harder as the technologies evolve.

## Pay-Per-Prompt?

A more flexible alternative might be to register the use of a 'copyrighted' entity when they are invoked in a text-prompt, so that the prompt '*Victorian cyberpunk diorama in the style of Greg Rutkowski*' would trigger a micro-charge, presumably with a prior warning to the end-user before execution.

The only two scenarios in which this would work would be in a web-facing API such as Stability.ai's [DreamStudio](#) (which uses the company's own Stable Diffusion), or OpenAI's [DALL-E 2](#) web interface (for which, unlike Stable Diffusion, there is no alternative means of access); or else a locally installed application that is far more locked down than open source distributions of Stable Diffusion (such as the hugely popular [AUTOMATIC1111](#) release), where the user experience would need to be *several orders of magnitude* better than available open source alternatives in order to gain traction, under the weight of such restrictions.





'Image of a courtroom scene, by Greg Rutkowski; insane detail, octane render' – the AUTOMATIC1111 distribution of Stable Diffusion, which runs in a standard web-browser connected to a Python runtime on a local installation of Windows (in this case).

As Bradford Newman observes, trying to solve the problem through record-keeping, either at the data-training or the generative stage, is fraught with challenges.

"So the issue is," he says, "will there be legislation requiring the humans behind algorithms to keep records of every search done, and what datasets were involved in each search? You can imagine how administratively burdensome and expensive that would get.

"You'd be storing an almost unimaginable amount of data for the possible eventuality that some litigant will later crawl out of the woodwork regarding the use of their copyrighted work in a dataset. It's almost undoable."

"But even that," he continues. "has problems, because I don't know what the right remuneration would be. If a source dataset has five million images, do all five million get an equal share? Is it only the input that matters, or is it only the output?"

"Your work may sit, trained into the dataset, and never contribute *anything* to any output from the system, by any user. Or it may represent a big contribution to the output, to the way that people actually summon up that data and make us of it.

"If we can't quantify the value of a particular contribution, then a scheme of that type could raise as many questions as it provides solutions. In terms of a legal remedy for the abstraction of copyrighted work, it would be a stop-gap measure at best."

Indeed, it seems logical to assume that cynical actors would eventually also seek to be included in datasets that compensate contributors equally for synthesized derived works – possibly with plagiarized or poor-quality work that's just good enough to not be [filtered out by CLIP](#) during the data-gathering stage, and ensure equal-parity payment with the most notable (and far more talented) dataset contributors.

## China, and the Race to the Bottom

In regard to computer vision technologies in general, and image synthesis systems in particular, it is hard, though not impossible, for liberal democracies to compete with autocratic or authoritarian regimes; but in the case of exploiting public data for machine learning research, this does entail engaging in a race to the bottom, where the pace is almost entirely led by China – a country that, as we've seen, now dominates the daily computer vision listings at Arxiv and other research publication platforms

China can set in place [a long-term plan](#) to achieve a global lead in AI research, and follow up on it without imperiling its funding at periodic election events; can take [all the time in the world](#) to (slightly) improve its [rock-bottom reputation](#) for respecting intellectual property or citizens' [rights to privacy](#); and can surveil its citizens at an [unprecedented scale](#), and [exploit](#) this vast swathe of data exclusively, if it chooses.

Most of the [improvements](#) in China's regulation of intellectual property have come about through its need to operate in the west, and through other international pressures to reform its laggard IP laws and enforcement policies.

However, there are notable areas where export and compliance are not the primary consideration, such as the application of computer vision for [military](#), [industrial](#), [surveillance](#) and, to an extent, domestic medical research – all pursuits from which, in respect to computer vision technologies, China can obtain a unilateral advantage – and, if it chooses, share this advantage with 'friendly' countries that may be less allied to western interests.

"China is really interesting," says† Kevin Buckley, a California-based international patent attorney specializing in IP law related to AI. "They legislate, obviously, in a very different way to the US."

"If they need greater control over the legal liability of a dataset and any ML/AI synthesis systems derived from it, they've already got the copyright law to do that. They can simply require the original contributors to sign a waiver that allows their images to be incorporated into a dataset, and that will kill the issue right there."

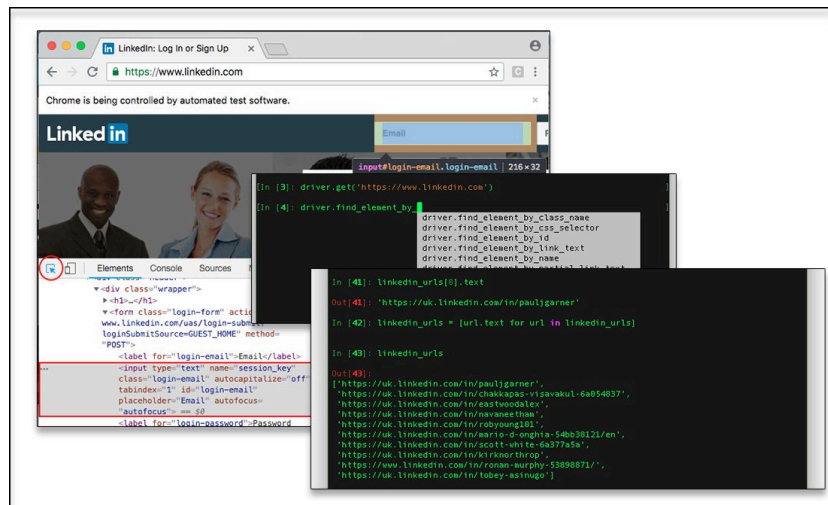
"Of course," he continues. "any western government, any government at all, can do that. But in a liberal democracy, that kind of decree is likely to encounter more friction and stir a greater amount of public debate."

## The Fair Use 'Lifeboat'

For these reasons, among others, western democracies seem likely to favor a *post facto* approach to the consequences of incorporating public data into image synthesis systems; but here too China has a massive advantage, since it can impose, with scant opposition, arbitrary laws that prohibit unsanctioned use of image synthesis technologies – indeed, in the case of deepfakes, it has [long since done so](#).

In the west, the potential enactment of such prohibitions is bound up in a web of often contradictory intra- and international regulations regarding the doctrine of [fair use](#). Therefore western governments neither have an easy way to *selectively* stem the use of web-scraped imagery for AI training, at least without tripping over multiple opposing policies and regulations (many of which might be imposed by partner states); nor any decisive and persistent powers to ban the fruits of the resultant synthesis systems – excepting that certain types of image content, synthesized or not, may already have been [prohibited](#) by state and/or national laws.

There are by now at least a couple of cherished prior art cases to which proponents of data-driven AI defer as dissenting voices rise: [Author's Guild v. Google](#), a [picaresque](#) legal journey that lasted ten years and ultimately preserved the search giant's right to trawl web-facing text data; and [HIQ Labs, Inc. V. LinkedIn](#), another 'general' web-scraping case in which LinkedIn's [sporadic](#) victories [did not ultimately](#) grant it legal protection from *ad hoc* data miners.



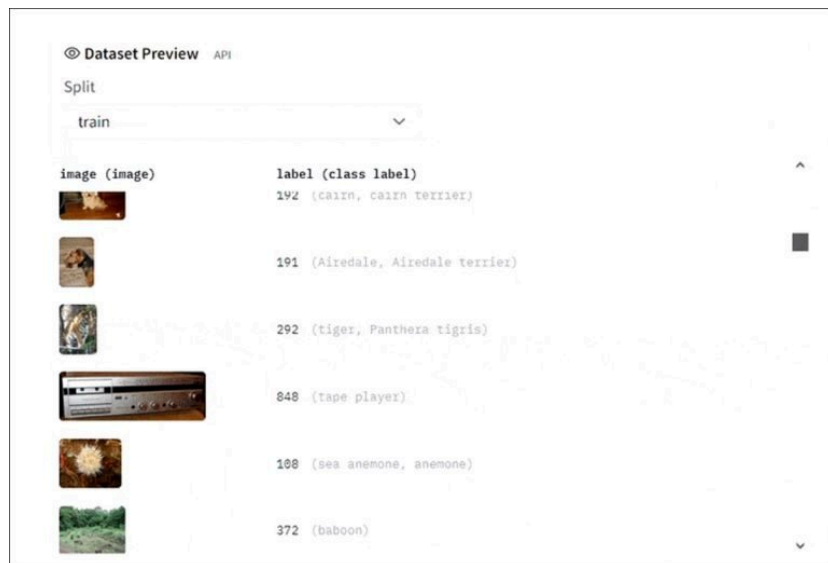
Images from a how-to article that demonstrates the creation of a webscraper for LinkedIn (ironically, hosted on LinkedIn). Source: <https://www.linkedin.com/pulse/how-easy-scraping-data-from-linkedin-profiles-david-craven> (archived: <https://archive.ph/KFOYf>)

Beyond these decisions (generally applicable only in the US), defenders of the *web>AI* pipeline cite whichever local or national enshrinements of 'fair use' seem likely to apply, or likely to best protect them.

However, the overturning of *Roe Vs. Wade* in June demonstrates that even stalwart and venerable legal rulings can yield to the passions and demands of the day, which makes commercial computer vision development a risky business.

## The ImageNet Conundrum

For instance, advice at [this discussion](#) at the PyTorch GitHub, about the legality of commercial projects that use the venerable and perhaps [overly-influential](#) [ImageNet](#) dataset and derivatives, exemplifies the fragile and provisional standing of ImageNet's commercial downstream descendants (though the same applies for several other 'foundational' image datasets, such as the [CELEBA](#) series).



Browsing the image data (and associated captions) in ImageNet. Source: <https://huggingface.co/datasets/imagenet-1k/viewer/default/train>

One poster at the thread noted that most of the [PyTorch TorchVision models](#), central fixtures in some of the most popular and well-diffused computer vision applications, are trained on an ImageNet subset that was [subsequently taken down](#); and that in any case, the [original applicable terms of use](#) prohibit commercialization.

Despite the fact that ImageNet's own terms, both historical [and current](#)<sup>\*</sup>, unambiguously ban the commercial use of the database, its [enormous influence and legacy](#) means that almost any current computer vision or image synthesis project, including hundreds or even thousands of

commercial applications, is in some way a downstream beneficiary of ImageNet – which is why, as with the above-mentioned GitHub conversation, the [default response](#) to questions about the legality of using it is to consult a good IP lawyer.

That's a fragile legal foundation for a business. If, for some reason, the commercial or general use of ImageNet were ever to become prohibited, either because it is bursting with the fruits of unauthorized and uncompensated web-scrapes (obtained under broadly interpretable fair use terms at the time), or simply because so many projects violate its clearly-stated terms of use, it's not unreasonable to say that the western computer vision sector would be subject to an enormous 'reset', and that a large swathe of commercial applications would be adversely affected, if not wiped out.

Further, it would not even be easily possible to start over with fresh data obtained under regulated licensing, since so many of the computer vision tools, libraries and frameworks that now operate on image data would themselves be part of a [poison legacy](#), having either been trained on ImageNet-based datasets, or regulated in training by ImageNet-benefited [loss functions](#), or whose performance metrics also owe some debt to ImageNet, possibly at several removes (i.e., a framework that uses a component that uses a library that uses an architecture that used ImageNet).

Therefore, any proposed western regulation addressing the non-authorized use of publicly available images for the purposes of training machine learning systems is going to have to give older datasets such as ImageNet some kind of safe harbor provision or statute of limitations if there is to be any hope of shortening China's lead in computer vision research in the near future.

By now, such 'ring-fencing' would almost certainly need to include the hyperscale image dataset [LAION5B](#) (which powers Stable Diffusion, among other initiatives), and its derivatives, with the need to extend this protection also to the [controversial Common Crawl](#) dataset from which the LAION series was derived.

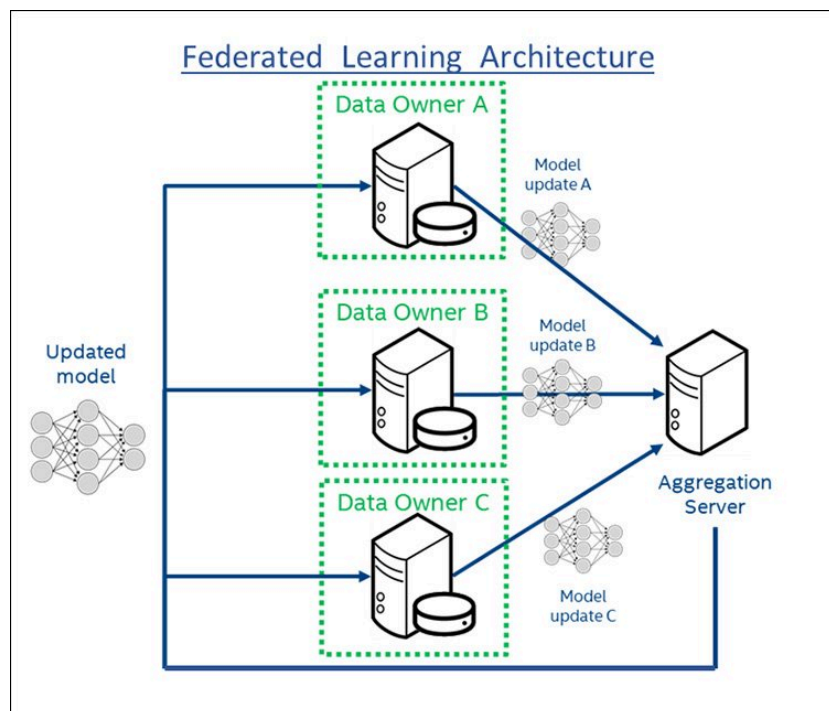
The scale of Common Crawl is staggering; at the time of writing, the latest Common Crawl scrape [comprises](#) 3.15 billion web pages, equivalent to 417TB (380 [tebibytes](#)) of uncompressed content. Realistically, any prohibition or limitation on web scraping would have to give amnesty to datasets that contain such copyrighted content – even extremely *recent* content, which many of the more vocal proponents of regulation might consider a tepid and unfair half-measure.

"I'm a wagering man," says Brad Newman. "and I would bet the house that whatever law we get around this will not be retroactive, but only valid from the day it's enacted, going forward."

But even having survived that legislative and PR hurdle, there remains the question of which system all this *ad hoc* 'image piracy' could possibly be replaced with, without ceding yet more ground to the Asian computer vision sector, which is in a better position to impose long-term legal exceptions, and to keep the data flowing.

## Federated Learning

One possibility is [federated learning](#), where the system being trained receives only the [weights](#) from privileged data.



Under federated learning, the original image data remains at the source, processed locally into weights, which are then passed through to a central training architecture, usually on another network. Source: <https://blog.openmined.org/untitled-3/>

Discarding the pixel data that yielded the weights is, in any case, standard procedure for data training pipelines; so in theory this is one approach that could allow for data sharing without the aggregation and dissemination of a massive corpora containing copyrighted material.

Though federated learning has proved an enticing prospect, particularly for [HIPAA](#)-constrained architectures that [must ensure patient privacy](#), it comes with its own set of risks. For instance, without liberal access to the source data, data poisoning attempts may be harder to identify, and the results may be less transparent, reproducible and accountable for aspects other than the privacy of the contributing entities.

Additionally, federated learning requires local machine learning resources of some kind. Though these are becoming [increasingly common](#) in mobile and smart devices, and otherwise low-complexity systems (such as domestic control systems and other 'edge' frameworks), their power consumption issues have yet to be entirely overcome – an [obstacle](#) to using medical federated learning systems routinely on smartphones, for instance.

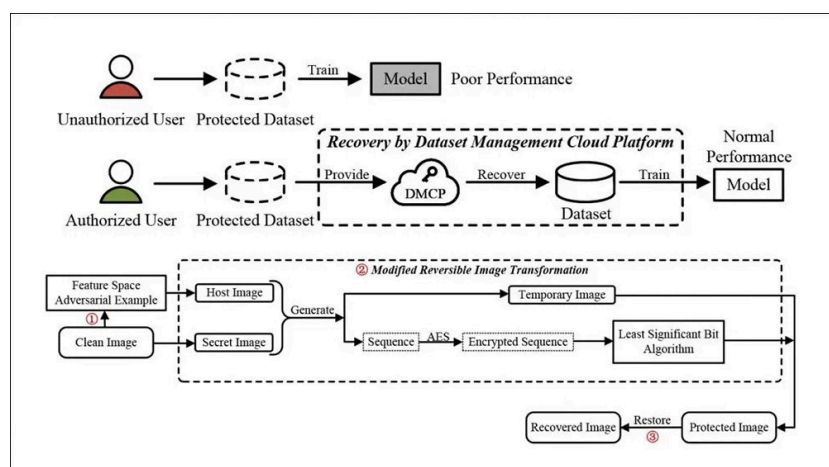
At greater scale, agreements could be wrought wherein host companies (such as LinkedIn) agree to provide weights from their content platforms – possibly through legislative constraint – with the local processing costs either compensated directly or through collateral benefits. This would make multi-terabyte/petabyte image collections a historical footnote.

The upside of federated learning is that it potentially provides a secure environment in which contributors may feel appeased by the system's supposed provision of anonymity, and be willing to explicitly share more (where not actually constrained to do so for the 'national good').

The potential downside, besides those already outlined, is that the training process may not remove all personally identifiable information, and that the process itself seems likely to be far harder to audit, in the event of controversy or complaint.

## Data Watermarking

Another approach that's gained traction in the last couple of years is [data watermarking](#), wherein the image data passing through the system is in some or other way encrypted or protected. One such method, a Chinese [initiative](#) from September of 2021, adds perturbations to images so that they are effectively useless for training in any 'unauthorized' machine learning pipeline.

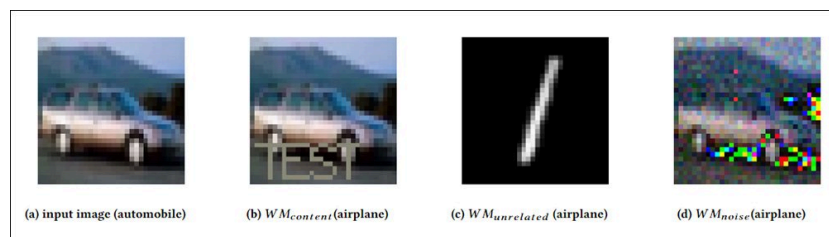


Unauthorized use of training images results in catastrophically degraded loss values in any unauthorized training schedule, under the 2021 system from the Nanjing University of Aeronautics and Astronautics. Source: <https://arxiv.org/pdf/2109.07921.pdf>

Though this method does not stop the process of web-scraping, or of data aggregation, it does ostensibly prevent such data from being used for any non-sanctioned purpose. Negatively, most of the schemes that have implemented this workflow have difficulty in preserving the quality of the images as the steganographic signals are added.

Additionally, this type of approach seems suited only for purpose-built datasets, probably not at the largest scale, since applying it to every image on the internet is an unlikely prospect, while the drop in human-perceptible quality would probably be unacceptable.

Variations on this methodology include a 2018 project to [watermark the output](#) of a generative model, though this is more concerned with model IP than user privacy or general IP protection; and a 2018 initiative from IBM, which sought to invest model data with patterns designed to [provoke miscaptioned images](#) for non-authorized frameworks, which would ensure that user-contributed data would only be exploited by a permitted model system, since manual re-captioning of vast swathes of data is impracticable.



Under the IBM system, embedded watermarks effectively constitute an attack on unauthorized deep learning systems attempting to train on the 'altered' data, by forcing them to misrecognize (and miscaption) images. Source: [https://gzs715.github.io/pubs/WATERMARK\\_ASIACCS18.pdf](https://gzs715.github.io/pubs/WATERMARK_ASIACCS18.pdf)

In most cases, for such schemes, the watermarking comes not only at the cost of reduced human-perceived image quality, but usually in some or other drop in loss accuracy during training.

## Data Laundering

The growing grass-roots movement to make AI accountable (even chargeable) for the images that generative systems train on is chiefly a speculation of artists and writers, along with a tranche of IP lawyers, in some cases perhaps propelled more by opportunism than idealism.

A greater political challenge than this may eventually come from more powerful and influential legislative voices, now beginning to hear concern that the research sector's *laissez faire* permissions to ingest and train public data (irrespective of terms set out by host domains or the platforms that run on them) constitute tacit [data laundering](#) – a breach of good faith of the original 'academic' terms under which valuable data was obtained prior to being monetized at scale.

As far as the US is concerned, this appears to be accepted as a type of 'trickledown economics', where the former conventional routes by which research was monetized (patents, publication and deployment) have been overtaken by the less clear pathways through open source culture, and by the very long upstream ancestry of 'idealistic' open source projects, in respect to commercial projects which are now benefiting from them.

Since, finally, such research leads to increased national income, albeit by a disingenuous route, it would seem to serve the original ends of privately and publicly-funded research: monetizable innovation.

So on the one hand, western legislators must contend with a new wave of public ire about the 'appropriation' of web-found content into free and open source generative systems, where the latter may subsequently transform into a [billion dollar](#) software release (from which the original, 'unwilling' contributors currently receive nothing), along with the complex tax and regulatory implications of tacit data laundering.

On the other hand, they are well aware that the global computer vision sector will increasingly depend on abundant – and *freely available* – data; and that Asia in general (and China in particular) is currently in a better position to gather and exploit such data.

Inconveniently, these two election issues are pulling in the opposite direction. Meantime, implementing *post facto* solutions, such as the regulation of AI-generated images, or [requirements for safeguards](#) in image synthesis systems, is a more politically charged and legally ambiguous prospect in the west than the east – and in any case could not have a wider-reaching influence without international legal consensus, and some basic synchronization of laws across partner states and countries.

## GitHub Copilot Lawsuit

On October 17th, the first major class action suit was [launched](#) on behalf of a 'proposed class' of millions of millions of GitHub users, with a [class action complaint](#) filed against Microsoft GitHub, for the use of open source code for the training of Copilot, a machine learning system that assimilates code from the popular repository site and provides AI-based coding services; OpenAI, Inc., for its use of source images to train DALL-E 2, and a handful of other related services.

One of the central named litigants, [Matthew Butterick](#), a lawyer and programmer working in concert from attorneys from the Joseph Saveri law firm, told us in an email of November 5th 2022:

*'I'm a lawyer, but for the last 10 years I've primarily worked as a creator and programmer. When I looked into Copilot, I came to the same conclusion that many open-source programmers already had — that it's built on widespread violation of open-source licenses.'*

*'As a technologist, I'm genuinely enthusiastic and optimistic about AI systems. But they're not exempt from the law. AI has to be fair [and] ethical for everyone. Companies making AI products have to remain accountable for their choices. It's disappointing but foreseeable that as more money flows into AI, the incentives are shifting, and so are the priorities.'*

*'At first we heard a lot about the existential threat to humanity posed by AI, and the importance of ethics, safety, and alignment. Those guardrails are already being removed, piece by piece. (Elon Musk was one of the earliest voices of caution, and yet he fired Twitter's AI ethicists this week.)'*

*'I teamed up with Joe Saveri and his firm to challenge GitHub Copilot because it's unfair and it represents a dire risk to open source. But I also hope to see a growing public conversation about how we the people want to situate AI in human society.'*

Commenting on the potential future of this new class-action suit, Bradford Newman told us:

*'The honest answer is that nobody knows what the outcome of this court battle will be. Both sides have strong arguments. Ultimately, the core "fair use v. infringement" issue for this AI use-case will wind up being decided by the appellate (or the United States Supreme) court unless the parties settle and/or Congress passes new laws.'*

## Open Source Matters: Copyright Law Vs. Contract Law

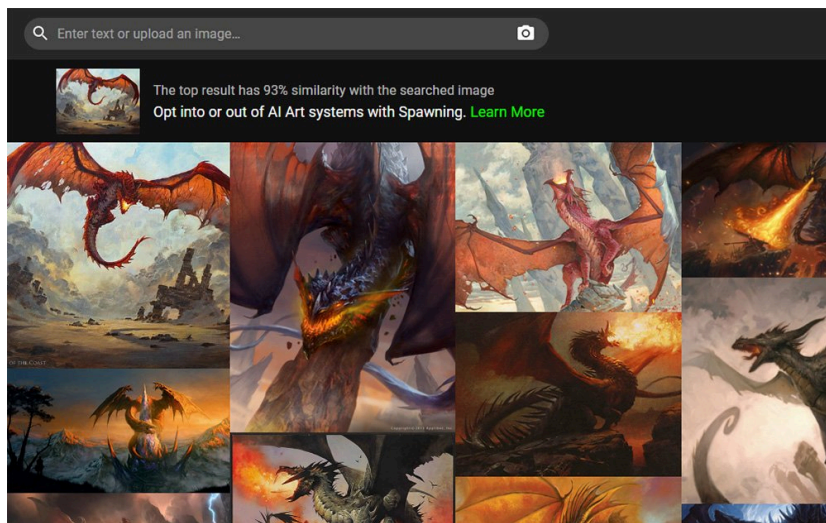
Though Kevin Buckley believes that new case law is needed to clarify the legal ambiguity of copyrighted material in AI-facing datasets, and that the Supreme Court will ultimately settle the matter, he contends that most of the relevant law has been on the books for some time.

"In the US at least," he says. "and broadly across the world, most of the issues around this are covered by our adoption in 1989~~†††~~ of Switzerland's Berne Convention. That law had existed worldwide since the 1880s, and was *not* drafted by the US government. So this is an area where we're already pretty closely aligned with the rest of the world."

"In any case, the copyright interest in a set of images is very, very weak. You could consider it in the same light as telephone directories, which compiled names, addresses and phone numbers into a collection, but which didn't result in works that could be legally protected from copying. Instead, they were protected by the logistical difficulty, at least at the time, of reproducing and distributing the collection.

"So if you gather up a whole load of web-scraped kittens into a dataset that gets trained into a model that can generate images of kittens, there is some limited copyright interest in that; but it's very weak."





The website [haveibeenentrained.com](https://haveibeenentrained.com) assesses the similarity between user's source images and known generated artworks, helping 'victims' of ad hoc dataset ingestion to know if their work has contributed to the output of generative systems. Source: [haveibeenentrained.com](https://haveibeenentrained.com)

Buckley believes that though current laws do not necessarily need to change in order to address copyright issues around AI's appetite for public data, the industry's treatment of open source licenses (such as the aforementioned ImageNet license) is a different matter.

"So the legal analysis is done," he says. "But in regard to open source, you're looking now at contract law, which is different in every jurisdiction.

"As lawyers, we try to de-risk our clients, who are using these datasets, and who are allowing others to use these datasets. And because the law is uncertain, we know that there's a weak copyright in these compilations. If somebody is utilizing that dataset, and our clients provide a pathway for that to happen, we de-risk it through contracts that relieve our clients from liability, so that the user of these datasets is the one who's exposed, legally."

## The Spice Must Flow

In the ordinary course of events, legislation would seek to address a public controversy *at source*, in some way seeking to limit access to requisite initial materials, or access to an enabling technical process. Though this is not guaranteed to work, it does give politicians an easy exit from public disapprobation, since they can claim to have addressed the issue to the best of their abilities, according to state-of-the-art remedies available (whether those remedies are truly effective or not).

In the case of the use of publicly-available image data for the training of new AI systems, the stakes are arguably too high to obstruct the academic and research sectors (who, as we've seen, are far from distinct from each other) from retaining image data access comparable to China.

This leaves *post facto* legislation as the only reasonable alternative. In the case of image synthesis systems, this could only prove effective by mandating the restriction and sanitized output of such systems through DALL-E 2-style, locked-down API access.

At the moment, this would make no difference, since the very powerful Stable Diffusion is already at large in the private sector and the hobbyist community, with the possibility of later systems, even from other companies than Stability.ai, being released to open source.

Eventually, however, the greater resources of the private sector would likely overtake and render outmoded the forks and development cycles of the hobbyist or enthusiast community, while larger players would become obliged to restrict access to newer and better models that greatly outstrip the 'amateur' branches in quality, driving interest towards them.

## Face Value

"I think there's an intentionality behind what China does," says Bradford Newman. "that is organized from the top down, based on the set of values that governs their society. We certainly don't have that intentionality; ours, on the surface, centers on how we can use AI to improve lives and make money in a capitalist system.

"I think the main question may be 'are we going to wind up unintentionally in the same place that other societies have wound up intentionally'?

"I don't think the answer to that is 'yes'; but there are caveats: I think we are going to give up a lot of individual privacy and liberty, but do so in a "willing" fashion. Americans love convenience, like everybody else. We love ease. And we're willing to give up a lot, and not ask too many questions.

"But we're not necessarily giving it up to a centralized governmental authority. We're giving it up to certain companies that are going to control the large datasets, who could potentially become like quasi-sovereign entities.

"What will they do with all this data? And should the data be more widely available, or curtailed by the courts? These are questions that Congress and the Supreme Court of the United States are going to have to address, and that's coming.

"The fundamental concern is, do we, as a western liberal democracy, understand the Faustian bargain we are making with what AI offers? The upside is diagnosing cancer in advance, and all of these healthcare applications that are societally positive. But there's also AI use cases for predicting crime, for recruitment, and many other areas that I don't feel so comfortable about.

*"And so I have concerns from a civil liberties standpoint; but I don't think it's the US government driving it – and I think that's a key distinction."*

*\* Archive snapshot taken by the author Wednesday, October 26, 2022 15:44:56*

*† In a conversation with us on 27th October 2022*

*†† <https://www.techdirt.com/2006/02/27/when-70-of-the-cost-of-blank-cds-goes-straight-to-the-recording-industry/>*

*††† Following the [decision](#) by Congress in 1976 to adopt the Berne Convention, it was finally [implemented](#) on March 1st 1989.*