

Disillusion Is Unemployed. The Illusion Is Hiring.

A sourced analysis of why calling every strange image “AI” erases investigation and awards machines human labour they did not perform.

2026.07.14 18:58 DENNIS HEDEGREEN ANALYSE V1.0

[HTTPS://HEDEGREENRESEARCH.COM/ARTICLES/DISILLUSION-IS-UNEMPLOYED/](https://hedegreenresearch.com/articles/disillusion-is-unemployed/)

Captain Disillusion spent years teaching the internet how not to be fooled by video. Then the internet discovered a faster method.

It writes **AI** underneath strange images it cannot immediately explain.

A creature moves strangely. AI. A camera travels somewhere a camera should not fit. AI. A polished product animation appears. AI. A painter’s textures look too smooth, a face looks too symmetrical, a photograph feels too dramatic, a practical effect is too good: AI.

The answer may occasionally be correct. But a correct guess is not an investigation. It does not tell us whether the image was generated, composited, animated, staged, simulated, edited, photographed under unusual conditions, or built by people who spent years learning how to make impossible things appear possible.

Captain Disillusion—Alan Melikdjanian behind the silver face paint—did not build his work around the word *fake*. His work began where that word stopped being useful.

What kind of fake? Which part of the frame was altered? Could the camera movement exist? Was the apparent miracle produced through tracking, masking, perspective, animation, practical construction, selective editing, or simply a real event that looked unfamiliar?

That distinction was the job.

The title of this article is therefore not a report about Melikdjanian’s career. In April 2026, Nebula announced a first-look deal with him and added his catalogue to the platform. The unemployment is rhetorical. It is the investigative method that parts of contemporary internet culture behave as though they no longer need.

Disillusion is unemployed.

The illusion is hiring.

“AI” is not a production history

Synthetic images did not begin with generative artificial intelligence. Artists were painting backgrounds, building miniatures, combining exposures, animating frame by frame and compositing visual layers long before diffusion models and text prompts.

Even **CGI** is not one method. A shot may combine hand-modelled geometry, photographed textures, simulation, motion capture, key-frame animation and weeks of compositing. A modern image may use machine learning without being generated from a prompt, or contain one generated element inside a production otherwise controlled by people.

The old binary—real or fake—was already too crude. The new binary—human or AI—is worse.

There are at least four separate questions:

1. **What does the image appear to depict?**
2. **Did the depicted event actually happen?**
3. **How was the image produced or altered?**
4. **Which people and systems performed which parts of the work?**

“AI” may offer a partial answer to the third question. It answers none of the others by itself.

A genuine recording can depict a staged event. A real event can be edited deceptively. A completely fabricated scene can carry an accurate production history.

As production methods expand, public language is contracting. One label is swallowing the history of visual effects and several distinct questions about evidence.

That creates at least three different failures:

1. Synthetic media can be mistaken for a recording of reality.
2. Real recordings can be dismissed as synthetic media.
3. Human-made CGI, animation, editing, photography and practical effects can be misattributed to generative AI.

The third failure is not merely a technical classification error. It changes who receives credit for the work.

When a human-built image is labelled AI without evidence, the machine is awarded labour it did not perform.

The Monet that became AI when someone said it was

On 12 May 2026, the conceptual artist SHL0MS posted a cropped reproduction of Claude Monet's *Water Lilies* (about 1915), held by the Neue Pinakothek in Munich, and presented it as AI-generated. Viewers were asked to explain why it was inferior to a real Monet.

Some visible respondents accepted the premise and supplied the defects.

They found the expected absence of intention. They criticised the composition, brushwork, depth, colour and texture as evidence of machine production. Others recognised the painting or challenged the setup. The post accumulated millions of views before the reveal circulated widely, but views are not a count of people who believed the label.

The machine had not painted the picture.



Claude Monet, *Water Lilies*, about 1915, Bayerische Staatsgemäldesammlungen – Neue Pinakothek München, inventory 14562. Reproduction: [Bayerische Staatsgemäldesammlungen](https://www.sammlung.pinakothek.de/en/artwork/6kLaEKXL8V) (<https://www.sammlung.pinakothek.de/en/artwork/6kLaEKXL8V>), CC BY-SA 4.0 (<https://creativecommons.org/licenses/by-sa/4.0/>).

The label had altered the viewing.

This was more interesting than a simple internet prank because the commenters did not merely repeat the statement that the painting was AI-generated. They looked at the image and discovered visual reasons why the statement must be true. The conclusion arrived first. Perception assembled the supporting evidence afterwards.

Research on AI labels points in the same direction. In a 2026 behavioural and eye-tracking study, participants gave higher evaluations to paintings labelled as human-created even when they could not reliably identify the actual source. The broader eye-tracking findings were mixed.

That reaction is not automatically foolish. The origin of a work can legitimately matter: art can be evidence of another person's attention, choices and experience.

But valuing human authorship is not the same as possessing a reliable detector for it.

Disliking generative AI does not make a person able to recognise it. A moral position can guide what we value. It cannot substitute for evidence about production.

The Monet episode exposed a basic vulnerability: once a category is supplied, viewers can begin seeing the category's expected features. "AI" does not merely describe the image. It can instruct people how to perceive it.

The audience is part of the detection system

The European Union's SOLARIS project studied 1,124 participants in Italy, Slovenia and the United Kingdom. Correct detection fell from about 60% to 37% for climate deepfakes and from 44% to 31% for immigration deepfakes when video quality moved from low to high.

The viewer still mattered. Participants with more positive attitudes toward the person shown were less likely to identify the manipulation. Higher media literacy predicted better detection.

The European Commission later described this as political alignment. The measure was narrower: attitude toward the depicted person, not a general match between the message and a left-right ideology. The detailed analysis is a preprint, and the project deliverable carrying the results was marked as not yet approved by the Commission.

The technical quality of the fabrication was only one variable.

The viewer brought another.

People do not usually sit down and consciously choose a lie over reality. The process is quieter: information that fits an expectation meets less resistance, while threatening information attracts more scrutiny. The evidentiary threshold moves without announcing that it has moved.

A politically convenient video may therefore meet less resistance, while an inconvenient video may attract more scrutiny. But the SOLARIS result does not establish that one person routinely makes both errors, or that ideology alone determines detection.

The information environment nevertheless permits both rules:

This fits what I expect, so the video is evidence.

This threatens what I expect, so the video is AI.

Reality demands that we sometimes revise ourselves. A private reality only demands that we revise the evidence.

This is why the deepfake problem cannot be solved solely by making people more suspicious. Suspicion is not neutral. A person trained only to distrust images may become harder to fool with one fabrication and easier to recruit into rejecting authentic evidence.

The result is not universal disbelief.

It is selective reality.

The artist who had to open Blender

The human-labour problem is not hypothetical.

On 12 September 2024, the 3D artist Alexandr SubSensus posted a short surreal animation to Reddit. Its finish was polished, its imagery unusual, and its visual style resembled qualities that audiences had begun associating with generated media. Commenters accused him of using AI.

SubSensus responded in the original thread with a viewport image. He and subsequent reporting identified a Blender, Cinema 4D and Octane Render workflow.

The argument did not simply disappear. The viewport produced more debate.

That detail matters. The artist had already made the work. He then had to perform a second job: reconstruct the evidence that he had made it.

The viewport is creator-supplied workflow evidence, not an independent audit of every asset and tool. It does establish the sequence relevant here: accusation, demand for process evidence, response, continued dispute.

The accusation did not only challenge authenticity. It reassigned authorship. Years of learning modelling, lighting, animation and rendering were compressed into an imagined prompt.

The artist as forensic archivist

In April 2025, the indie developer Stamina Zero released a trailer for *Little Droid* through PlayStation's YouTube channel. Its thumbnail showed a small robot in a glossy visual style. Commenters declared that the game had been "ruined" by AI art.

The developer said the cover had been commissioned from an artist named Olga Kochetkova and linked her portfolio. The accusations continued.

The studio asked Kochetkova for intermediate sketches and the layered Photoshop file, then assembled a 65-second process video as damage control. It shows three monochrome stages, the finished illustration, a substantial layer stack and elements being toggled in Photoshop. Some viewers accepted it. Others objected that it did not show every mark being drawn.

The video is evidence of process, not a complete audit, and the source PSD was not independently inspected. But the public sequence is clear: accusation, retrospective evidence, then a demand for evidence that would have required recording the work from its beginning.

The artist is no longer merely asked to submit a work.

The artist may be expected to preserve a forensic archive of their own innocence.

Show the layers. Show the sketches. Show the timestamps. Show a screen recording. Show the camera original. Show the export history. Show that no undisclosed model touched the process at any point.

For competitions, publishers or clients, requesting relevant process evidence can be reasonable. But when that logic becomes a general presumption, process documentation turns into reputational insurance.

The burden will not fall evenly. Old work may have no surviving project files. A painter may not film every canvas. A studio may have confidential workflows. A disabled or hybrid artist may use assistance tools that do not fit a clean human-versus-machine category.

The pressure can also reshape aesthetics. Stamina Zero's developer reported receiving advice to avoid art styles resembling generated images. But generative models learned those aesthetics from existing human work. The result is absurd: a machine imitates artists, and then artists are warned not to resemble the machine.

Evidence can always be declared insufficient. A process video can be called staged. Layers can be called fabricated. Similarities can be found after the conclusion has already been chosen.

There is a difference between asking for evidence and constructing an accusation that no evidence is permitted to defeat.

When the defenders of artists erase the artist

Opposition to generative AI is not an imaginary problem invented by technology companies. Artists have raised serious questions about training data, consent, compensation, employment and market power. Careless accusations should not be used to dismiss those interests.

But a movement can defend a legitimate principle through an illegitimate inference.

An AI accusation can make a factual claim about production and deliver a moral verdict at the same time. Once "AI" means theft, laziness, fraud, environmental harm and job destruction, identifying it feels less like attribution and more like exposing wrongdoing. That raises the emotional reward for certainty.

A viewer writes "AI slop" beneath a human-made animation and may sincerely believe they are defending human creativity against automation.

In that moment, they have done the opposite.

They have removed the human from the work and transferred the achievement to the machine.

This is the **anti-AI omnipotence paradox**: in trying to expose artificial intelligence everywhere, critics can help manufacture the impression that it is already capable of everything.

AI boosters and reflexive AI opponents appear to occupy opposite cultural positions, but they can converge on the same mythology.

The booster says:

The model can already replace the artist.

The reflexive opponent looks at an artist's work and says:

The model must have made this.

Both accounts inflate the autonomy of the system. Both compress the human contribution into a prompt—or invent a prompt where none existed. The artist disappears twice: first into the training narrative, then into the accusation.

There is not yet a reliable dataset establishing which political or cultural group produces the most false AI attributions. That question should be tested rather than assumed.

The documented cases show that false attribution can force creators to spend time producing proof. Repeated often enough, it would also reinforce the booster's market story: the human work has already disappeared. This article has not measured how often that happens or which group does it most.

The contradiction does not require a prevalence claim. Wherever a defence of human craft assigns human work to a machine without evidence, the defence has turned into erasure.

The liar's dividend

The attribution crisis does not stop with art.

Convincing synthetic media creates a benefit for people caught by authentic evidence. Researchers call it the **liar's dividend**: once fabrication is plausible, a politician or other powerful actor can dismiss genuine damaging material as misinformation or a deepfake.

Survey experiments published in the *American Political Science Review* found that false misinformation claims could help politicians recover support after text-based scandal reports. Effects were more limited for audiovisual evidence. Even so, the study shows how uncertainty can become useful: a denial need not prove that a report is fake if it can make the truth feel unknowable.

That is the deeper danger of treating "AI" as a complete answer.

The label can help a false image enter reality, help authentic evidence escape accountability, or make human craft appear machine-made. The content changes. The attribution move is similar.

Provenance, not vibes

YouTube's disclosure policy is more precise than the average comments section. It distinguishes realistic generated or meaningfully altered content from minor edits and production assistance. Not every digital intervention is the same act.

C2PA Content Credentials can preserve signed information about an asset's origin and edits. But provenance is not a truth machine. A credentialed recording can depict a staged event; a genuine photograph can carry a false caption; an honest image can lose its metadata as it travels. An April 2026 independent security preprint went further, arguing that the current specifications should not yet be relied upon for high-stakes journalism or legal evidence.

Absence of a credential cannot become evidence of guilt.

The SOLARIS project's work with journalists points towards the more durable method. In a simulated exercise with ANSA, journalists described a verification hierarchy: check sources first, test the context second, and use technical anomalies third.

That ordering matters. A logo can be checked against the broadcaster's site. A speech can be checked against other people present. A political claim can be tested against the institution and the date. Only then come the unnatural voice or mismatched lips.

The journalists knew in advance that some material was fabricated, and the follow-up focus group involved three people. The exercise documents a method, not a general detection rate. It also recorded a useful failure: an authentic, decade-old Trump video was initially treated as fake because it did not fit his current position. Context must include time, or it becomes another vibe.

That is slower than writing "AI."

It is also the work.

A responsible visual inquiry should ask:

- What exactly is being claimed?
- What is the earliest available source?
- Is the production history documented?
- Which parts are captured, staged, edited, generated or unknown?
- Does independent evidence confirm the depicted event?
- What level of certainty is justified?

The answer may be **AI-generated, traditional VFX, edited but real, staged, a hybrid workflow—or we do not know.**

"Unknown" is not a failure of media literacy.

Unsupported certainty is.

The illusion is hiring

Captain Disillusion's method was never valuable because he could name software from a few pixels. It was valuable because he treated a visual claim as something that could be investigated. He separated categories, reconstructed mechanisms and left room for uncertainty.

Generative AI does not make that work obsolete. It multiplies the plausible production histories and makes the method more necessary.

The danger is not only that machines can generate convincing images. It is that **AI** becomes the name we give an image when we no longer want to understand it.

That shortcut can fool us in every direction. It can conceal a fabrication, hide an older form of manipulation, release authentic evidence from accountability, or erase the people whose craft produced the image.

We abandoned the job he taught us to do.

Sources

1. Nebula, “Captain Disillusion Inks First Look Deal with Nebula Studios”
(<https://press.nebula.tv/captain-disillusion-inks-first-look-deal-with-nebula-studios/>), 3 April 2026.
2. Susan Gerbic, “The Man Behind the Makeup: An Interview with Captain Disillusion”
(<https://skepticalinquirer.org/exclusive/the-man-behind-the-makeup-an-interview-with-captain-disillusion/>), *Skeptical Inquirer*, 18 July 2016.
3. SOLARIS, “Use cases co-creative evaluation”
(https://projects.illc.uva.nl/solaris/uploaded_files/inlineitem/D5.2_final-tobesubmitted_with_EC_disclaimer.pdf), Deliverable D5.2, delivered 30 October 2025; document marked submitted but not yet approved by the European Commission.
4. Nejc Plohl et al., “How deepfake quality, media literacy, and personal attitudes shape detection, liking, and social media sharing of political deepfakes” (https://doi.org/10.31234/osf.io/knvby_v1), PsyArXiv, 19 May 2025.
5. European Commission, “The deepfake dilemma: can synthetic media be used for good?”
(<https://projects.research-and-innovation.ec.europa.eu/en/projects/success-stories/all/deepfake-dilemma-can-synthetic-media-be-used-good>), SOLARIS project success story, 9 July 2026.
6. SHL0MS's original post (<https://x.com/SHL0MS/status/2054280631807316329>), 12 May 2026; Bayerische Staatsgemäldesammlungen, Claude Monet, *Water Lilies*, about 1915
(<https://www.sammlung.pinakothek.de/en/artwork/6kLaEKXL8V>), inventory 14562; Hannah Vanek, “In Conversation With SHL0MS” (<https://opensea.io/blog/articles/in-conversation-with-shl0ms>), OpenSea, 19 May 2026.
7. “When Labels Matter More: Behavioral and Eye-Tracking Evidence on Aesthetic Judgment of AI-Generated and Human-Created Paintings by Lay Viewers”
(<https://journals.sagepub.com/doi/10.1177/02762374261441560>), *Empirical Studies of the Arts*, published online 13 April 2026.
8. Alexandr SubSensus, “My new blender Artwork”
(https://www.reddit.com/r/blender/comments/1fey5uy/my_new_blender_artwork/), Reddit/r/blender, 12 September 2024.
9. Joe Foley, “Blender artist has to prove their work isn’t AI”
(<https://www.creativebloq.com/art/digital-art/blender-artist-has-to-prove-their-work-isnt-ai>), *Creative Bloq*, 13 September 2024.
10. Lana Ro, “It’s definitely AI!”
(https://www.reddit.com/r/gamedev/comments/1jwbe09/its_definitely_ai/), Reddit/r/gamedev, 10 April 2025.
11. Stamina Zero, “The stages of drawing the cover art for the Little Droid game”
(<https://www.youtube.com/watch?v=QZFZ0YTxJEk>), YouTube, 10 April 2025.

12. Nicole Carpenter, “A real issue: video game developers are being accused of using AI—even when they aren’t” (<https://www.theguardian.com/games/2025/jun/26/video-game-developers-using-ai-even-when-they-arent-stamina-zero>), *The Guardian*, 26 June 2025.
13. Kaylyn Jackson Schiff, Daniel S. Schiff and Natália S. Bueno, “The Liar’s Dividend: Can Politicians Claim Misinformation to Evade Accountability?” (<https://doi.org/10.1017/S0003055423001454>), *American Political Science Review*, published online 20 February 2024; print volume 119(1), 2025.
14. YouTube Help, “Disclosing use of GenAI content” (<https://support.google.com/youtube/answer/14328491?hl=en-GB>), accessed 14 July 2026.
15. C2PA, “C2PA and Content Credentials Explainer” (<https://spec.c2pa.org/specifications/specifications/2.4/explainer/Explainer.html>), specification version 2.4, accessed 14 July 2026.
16. Enis Golaszewski et al., “Verifying Provenance of Digital Media: Why the C2PA Specifications Fall Short” (<https://arxiv.org/abs/2604.24890>), arXiv preprint, 27 April 2026.
17. Jocelyn Noveck and Matt O'Brien, “Visual artists fight back against AI companies for repurposing their work” (<https://apnews.com/article/7ebcb6e6ddca3f165a3065c70ce85904>), Associated Press, 31 August 2023.

RELATION MEMORY

EXTENDS [Recognition Is Not Verification](#) · Moves from the limits of recognition to the attribution error that occurs when a visual category is supplied before the evidence.

NEARBY [Surface Is Not Verification](#) · Both pieces argue that a convincing surface cannot carry the burden of a verified production history.

CONTINUES [The Tool Learned to Say Not Yet](#) · The same refusal applies to visual attribution: unknown is a stronger answer than unsupported certainty.