
What AI Asks of Open Access

By ALISON MUDDITT | MAR 31, 2026

ARTIFICIAL INTELLIGENCE | BUSINESS MODELS | INFRASTRUCTURE | INNOVATION | OPEN ACCESS | OPEN SCIENCE
| RESEARCH INTEGRITY

Two pieces landed in my reading queue recently that, taken together, crystallized something I've been trying to articulate for a while. The first (<https://www.authorsalliance.org/2026/03/16/library-and-archives-101-ai-and-the-false-promise-of-control/>), by Dave Hansen of the Authors' Alliance, makes a careful and important argument about why libraries and archives are making a mistake when they respond to AI by asserting institutional control over their collections. The second, (<https://scholarlykitchen.sspnet.org/2026/03/11/what-publishing-leaders-say-about-ai-when-theyre-not-on-panels/>) by Todd Toler and Angela Cochran, is a rich, reported account of where scholarly publishers actually stand with AI – the deals being done, the strategies forming, the organizational gaps widening. Both pieces get a great deal right but neither quite fits where PLOS sits. Working out why has helped me think more clearly about what AI actually asks of a mission-driven open science organization like PLOS.

Dave's core argument is that treating collections as assets to be leveraged, conditioning access on institutional control, demanding compensation – is not just strategically self-defeating but a values failure. And his most useful point: the organizations that navigated technology transitions well, including PLOS, did so by building infrastructure that offered *more* access, not less. Our founding was an act of construction, not restriction.

Todd and Angela's piece converges on a similar insight from a different direction. The most defensible publisher assets in an AI economy aren't the content corpus but the authority: the capacity to evaluate, certify, and set standards. APA Labs (<https://www.apa-labs.com/>), which evaluates AI mental health tools against psychological science standards rather than licensing content to build them, is their most striking example. The content can be compressed into model weights or assembled into someone else's context layer but the judgment cannot.



Both of these insights matter for PLOS, but here's where the framing starts to break down. Todd and Angela's piece pushes us past what are already the "wrong" questions about AI. But it still feels fundamentally organized around a monetization question: how do publishers participate in AI economics without giving away the store? That's a legitimate and urgent question if your content is behind a paywall. For PLOS, it's the wrong starting point: our content is CC-licensed and always has been open. Because of that, PLOS was among the earliest and largest sources of open, peer-reviewed scientific content available for AI training — not because anyone failed to stop it, but because that's precisely what CC licensing is for.

When PLOS content was cited as an example of openly available scientific literature that had made its way into AI training data, my reaction was relief, along with a sharper sense of responsibility. I'd far rather AI systems reason from a substrate shaped by authoritative, peer-reviewed science than from the vast ocean of unvetted information that makes up most of the internet. But that preference only means something if the content going in is structured, current, and integrity-preserving – which is exactly what the rest of this piece is about.

There's much that resonates with me in Dave's piece, and in this clear assertion, I could easily substitute "libraries and archives" for "PLOS":

...ing control, demanding compensation, and conditioning access on the institution's ability to dictate the terms of downstream use – that is not what libraries and archives are for."

But he's writing from a stewardship perspective — libraries as custodians of others' materials, navigating between donor obligations, community relationships, and access commitments. PLOS is a publisher and as such, we don't just preserve and provide access to the scientific record; we actively shape what enters it. That's a different kind of responsibility, and it points toward a different set of questions.

So, what are the right questions for a mission-driven open access publisher in an AI environment? If we're honest about what open access was supposed to be for, then open access and AI are not naturally in tension. For PLOS, open access was never the goal, it was the means to the end. The goal, our founding question, the one we're returning to in our current strategic work, was accelerating science: getting research findings into the hands of researchers, clinicians, policymakers, and the public as quickly and freely as possible, so that knowledge could build on knowledge without artificial friction. For twenty-five years, that meant fighting the paywall. The question AI forces us to ask is whether openness continues to serve that goal, and under what conditions. I'd argue it does but the conditions matter enormously, and they center on trust and integrity in ways that neither of the pieces I cited fully addresses.

The instinct running through much of the current industry discussion is to respond to the substrate quality problem by reasserting publisher control: prioritizing the Version of Record, conditioning AI access on licensing agreements, building frameworks that position publishers as necessary gatekeepers of quality. It's a coherent strategy if your content is behind a paywall and your business model depends on controlling access. But it mistakes the container for the thing itself. The quality signal that matters isn't who controls access to peer-reviewed content — it's whether the integrity markers that make that content trustworthy are open, structured, and machine-readable enough to actually function in an AI environment. Those are very different problems with very different solutions.

Here's the specific problem. PLOS content is especially available to AI systems, by design. Because of the scale of our corpus and the breadth of our CC licensing, that means we're disproportionately present in the substrate that AI systems are reasoning from. That creates both more exposure and more responsibility. AI systems cannot reliably distinguish a well-powered randomized controlled trial from an underpowered observational study, a retracted paper from a current one, a finding that has replicated robustly from one that hasn't. These failures aren't random, they're embedded at specific points in how these systems are built, from what goes into training data, to how quality signals are or aren't applied during tuning, to whether provenance is even surfaced to the user. Retrieval-augmented systems can do better, provided they are connected to current, well-structured sources and designed to use them, but this is a design choice, not a given.

Every publisher with a peer review process would say that process exists to make those distinctions. But not every publisher requires open data and open methods as a condition of publication, making those distinctions auditable by anyone, human or machine. Not every publisher has built the publication ethics infrastructure to keep the record honest as understanding evolves. And even where the values are shared, commercial publishing models create real pressures — on speed, on volume, on revenue per article — that can work against the kind of investment in integrity infrastructure that doesn't have an obvious short-term return. For PLOS, that particular tension doesn't exist in the same way

These aren't peripheral features of what PLOS does. They are, in Todd and Angela's terms, the authority that cannot be dissolved into weights — and “dissolved” is precisely the right word. Training-based systems compress patterns from content into model weights, where updates, corrections, and retractions don't automatically follow. And in an AI environment where the scientific literature is increasingly a substrate for automated reasoning rather than a collection of human-readable arguments, the integrity of that substrate matters more, not less.

Where does that leave us? PLOS's role in an AI world isn't to protect our content from AI, or to monetize AI access to it, or even primarily to build AI tools on top of it. Our role is to ensure that what PLOS publishes remains trustworthy infrastructure for scientific reasoning, whether that reasoning is done by humans or machines. That means doubling down on the things that make a PLOS paper mean something: rigorous (and increasingly open) peer review, mandatory open data, transparent methodology, robust correction and retraction processes. It means investing in the machine-readable signals (persistent identifiers, version histories, structured metadata, and clear licensing) that allow AI systems to do something they currently do poorly: reason about the reliability and currency of what they're drawing on. Open access is necessary but not sufficient here. Even CC-licensed content loses its attribution and provenance in most AI systems — the licensing that was designed to make reuse possible turns out to be largely illegible to the machines doing the reusing. For AI systems to use research responsibly, availability must be paired with structured, machine-readable signals that convey not just licensing terms but provenance, current status, and integrity.

None of this works at the scale AI demands without shared, open infrastructure and we have chronically underfunded the organizations building it. Crossref and Crossmark make it possible for a retraction to propagate across systems; ORCID links contributors to their outputs; Datacite makes data citable and traceable. If AI makes the integrity of the scientific record more consequential, our failure to resource the infrastructure that maintains that integrity becomes not just an embarrassment, but a genuine risk. Publishers, infrastructure providers, and AI developers need to work together on how publication ethics and research integrity norms are actually encoded in training data and product design, not left to chance or to whoever is building fastest.

The R&D work we've been doing over the past two years has pushed us toward a more active formulation of this. The most interesting question isn't how to make existing published articles more trustworthy for AI use, it's whether the article itself is still the right unit of scientific communication, and whether the infrastructure we've built around it is actually serving science or just serving a publishing system that science has grown accustomed to. I'll say more about this next month. For now: this work has deepened rather than complicated my conviction that openness and trustworthiness are not in tension — but that realizing both in an AI environment requires building something genuinely new, not just defending what we have.

The organizations Dave names as the open community's most consequential victories (HathiTrust, PubMed Central, arXiv, PLOS) all followed the same logic: how to the values and trust that access serves the mission better than control. AI doesn't change that logic, it intensifies it. The question isn't whether to open but whether those of us who believe in open science can build what it actually needs next — together, and with enough urgency to matter. We have the vision we have, I think, the vision. What we've been less good at is the collective action that turns both into infrastructure. AI won't wait for us to figure that out.



Alison Mudditt

Since June 2017 Alison has been Chief Executive Officer of the Public Library of Science (PLOS), an organization on a mission to drive open science forward with measurable, meaningful change in research publishing, policy, and practice. Prior to PLOS, Alison served as Director of the University of California Press and as Executive Vice President at Sage. Alison serves on the MIT Libraries Visiting Committee and the American Chemical Society's Governing Board for Publishing. A regular speaker at industry meetings, Alison also writes for the Scholarly Kitchen blog. Her more than 35 years in the publishing industry also include leadership positions at Taylor & Francis and Blackwell Publishers.

Discussion

3 THOUGHTS ON "WHAT AI ASKS OF OPEN ACCESS"



My understanding is that the CC license is not relevant for AI training and this falls under Fair Use of copyrighted material so is independent of the licensing of the content (though not its availability). We should also be wary of invoking CC BY as the justification rather than Fair Use because of the attribution clause: this is not practical for text mining, and in the case of current LLMs their very nature (probabilistic) means that attribution is not verifiable.

By RICHARD SEVER | APR 1, 2026, 5:47 AM



As always (and as we have noted repeatedly in these pages), it is important to differentiate between “attribution” (as specified by CC licenses) and citation (as is common practice in scholarly research). The attribution requirement of a CC license refers to reuse of the material, that is, republication of those words/images. For use of the material (reading, using it as the basis for further thought, or in this case, ingesting it into an LLM and having it play some role in influencing the answer that is spat out), attribution is not required. It's only when the material itself is reused in a manner where copyright would come into play (e.g., if the LLM spat out verbatim an entire book chapter). It is good scholarly practice to cite one's sources (this information came from here, this piece had an influence on the thinking here), but it is not mandatory, and that practice is beyond the scope of what a CC license requires.

By DAVID CROTTY | APR 1, 2026, 8:50 AM



Yes – the attribution vs citation distinction has been problematic from the outset: a license intended for creative works has always not quite suited the needs of academic research (as I recall, PLOS originally considered using CC0). And of course LLMs have prompted serious concerns about information provenance, credit and re-use – and a general lack of reciprocity by AI companies – from CC leadership and supporters. As I note above, the fact that when sources are provided, these may not reflect the true basis of the information generated makes it all the more problematic.

By RICHARD SEVER | APR 2, 2026, 7:04 AM



The mission of the Society for Scholarly Publishing (SSP) is to advance scholarly publishing and communication, and the professional development of its members through education, collaboration, and networking. SSP established The Scholarly Kitchen blog in February 2008 to keep SSP members and interested parties aware of new developments in publishing.

The Scholarly Kitchen is a moderated and independent blog. Opinions on The Scholarly Kitchen are those of the authors. They are not necessarily those held by the Society for Scholarly Publishing nor by their respective employers.

ISSN 2690-8085