

Guest Post – When AI Helps Write Research: What Happens to Lived Experience?

By SYLWIA FRANKOWSKA-TAKHARI | MAY 21, 2026

ARTIFICIAL INTELLIGENCE | DIVERSITY, EQUITY, INCLUSION, AND ACCESSIBILITY | ETHICS | EXPERIMENTATION
| INNOVATION | PRESERVATION | RESEARCH | RESEARCH INTEGRITY | SOCIAL ROLE | SOCIOLOGY | TECHNOLOGY

Editor's note: Today's post is by Sylwia Frankowska-Takhari, Accessibility Specialist for Springer Nature. Reviewer credit to Chef [Diannandra Roberts](https://scholarlykitchen.sspnet.org/author/diannandra-roberts/) (<https://scholarlykitchen.sspnet.org/author/diannandra-roberts/>).

Artificial intelligence is increasingly becoming part of the authoring process. AI-powered tools can now help researchers discover and synthesize literature, structure manuscripts, refine academic language, and generate preliminary thematic summaries.

Recent discussions in *The Scholarly Kitchen* (<https://scholarlykitchen.sspnet.org/>) have highlighted the opportunities and risks associated with this shift, including concerns about [research integrity](https://scholarlykitchen.sspnet.org/2026/03/09/guest-post-the-perils-of-using-generative-ai-to-perform-research-tasks-editors-and-publishers-viewpoints/), [bias](https://scholarlykitchen.sspnet.org/2026/03/09/guest-post-the-perils-of-using-generative-ai-to-perform-research-tasks-editors-and-publishers-viewpoints/) (<https://scholarlykitchen.sspnet.org/2026/03/09/guest-post-the-perils-of-using-generative-ai-to-perform-research-tasks-editors-and-publishers-viewpoints/>), and the changing role of [scholarly judgement and practice](https://scholarlykitchen.sspnet.org/2026/04/29/guest-post-when-thinking-is-outsourced-a-warning-from-a-scientist-trained-before-ai/) (<https://scholarlykitchen.sspnet.org/2026/04/29/guest-post-when-thinking-is-outsourced-a-warning-from-a-scientist-trained-before-ai/>). As AI writing assistants become embedded in scholarly workflows, another question emerges: what happens to [lived experience](https://www.nihr.ac.uk/different-experiences-framework-considering-who-might-be-involved-research) (<https://www.nihr.ac.uk/different-experiences-framework-considering-who-might-be-involved-research>) when its interpretation is mediated by AI?

To explore this, I used publicly available excerpts from two peer-reviewed qualitative studies about disabled academics. *Experiences of Researchers with Disabilities at Academic Institutions in the United States* (<https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0299612>) examines how researchers with disabilities navigate academic institutions, disclosure, support, and professional participation. *The Insider View: Tackling Disabling Barriers in Higher Education* (<https://link.springer.com/article/10.1007/s10734-019-00479-0>) explores disabled academics' experiences of disabling structures and cultures within higher education.



I pasted selected passages into Microsoft Copilot (<https://copilot.microsoft.com/>), a popular and widely used AI system (<https://techcrunch.com/2026/04/29/microsoft-says-it-has-over-20m-paid-copilot-users-and-they-really-are-using-it/>), and asked a simple question:

“Summarize the key themes across these excerpts.”

The resulting output was structured, coherent, and professionally written. It identified systemic barriers, administrative burdens, stigma, lack of institutional support, and cultural norms of overwork. It read like the kind of executive-ready synthesis institutions often prefer: clear, organized, and balanced.

What the system produced looked like a polished editorial summary, the kind of language commonly used in institutional reports or manuscript abstracts.

Nothing in the summary was factually incorrect. The themes were recognizable. And yet something subtle had shifted.

Where participants had described being “kicked out,” left “on your own,” or navigating “excruciating” health conditions, the AI summary translated these experiences into language such as “lack of support,” “emotional toll,” and “administrative barriers.” Where the original authors had critiqued structural ableism within academic institutions, the summary spoke more generally about “barriers” and “challenges.”

The tone of the narrative had changed. The shift becomes easier to see when comparing specific excerpts with the language produced by AI summarization.

EXAMPLE 1

Original testimony: “I was kicked out of my doctoral program when I was diagnosed with schizophrenia.”

AI synthesis: “Participants reported experiences of identity disruption and emotional strain.”

An account describing institutional exclusion becomes reframed primarily as a personal experience of disruption.

EXAMPLE 2

Original testimony: “I always thought there would be a system in place... but there isn’t. It’s very much you’re on your own.”

AI synthesis: “Participants described administrative barriers and the burden of self-advocacy.”

A narrative of institutional abandonment is translated into the language of procedural challenges.

EXAMPLE 3

Original testimony: “The whole problem with academia is it presumes that you’re able to put in a 60-hour week.”

AI synthesis: “Participants noted cultural norms of overwork and inflexibility.”

A critique of institutional labor expectations becomes framed as a descriptive cultural observation.

Across these examples, the themes remain broadly recognizable. The AI system identifies the relevant issues and groups them coherently. But *the framing shifts*.

Emotionally charged testimony becomes translated into neutral professional prose. Structural critique becomes administratively legible. Experiences expressed as exclusion, abandonment, or anger are reframed as challenges, barriers, or support gaps.

The themes are preserved. But *the moral force carried by the original testimony is softened*.

The experiment illustrates a subtle but important transformation in how testimony is represented. One way of describing this dynamic is **moral compression**: a process in which the themes of lived experience remain visible while the emotional intensity and structural critique carried by testimony become compressed through automated summarization.

Rather than introducing factual error, AI synthesis redistributes emphasis, subtly reshaping how lived experience is interpreted and communicated.

Compression and Redistribution

The shift is easy to overlook. Yet it highlights a dynamic that may become increasingly relevant as AI writing assistants enter scholarly publishing workflows.

AI systems are trained to produce coherent, balanced, and professionally legible prose. When summarizing qualitative material, they seem to translate emotionally charged testimony into neutral thematic language. Structural critique may become reframed as general challenges. Experiences of exclusion may appear as “barriers,” “support gaps,” or “administrative burdens.”

The result is not inaccurate, but *it redistributes emphasis*. Accounts that originally conveyed institutional harm or structural exclusion may become expressed in language that is more administratively legible and easier for organisations to absorb.

In qualitative research, that shift matters. Lived experience accounts do not simply describe problems. They reveal structural conditions. Emotional intensity: anger, frustration, exhaustion, often signals deeper institutional failures.

When such signals are softened through automated summarisation, the resulting synthesis may preserve themes while redistributing the moral weight carried by testimony.

The output remains coherent and informative. But *the texture of the narrative changes*.

Why Outlier Voices Matter

These dynamics are particularly relevant in accessibility and disability scholarship.

Disabled voices often appear as statistical outliers in datasets. Experiences of episodic illness, fluctuating health conditions, and institutional discrimination do not always cluster neatly into dominant patterns. A single narrative may reveal systemic barriers that affect relatively small groups of people but have severe consequences.

In such contexts, low frequency does not mean low importance.

Yet AI systems are optimized to detect patterns and prioritize recurrence. Repeated experiences become central themes, while less common experiences may be treated as peripheral.

When authors rely on AI tools to summarize literature, organize qualitative material, or refine manuscript language, those tools may inadvertently smooth the sharp edges of testimony, particularly testimony that challenges institutional norms.

The risk is not that AI tools will misrepresent disability experiences. *The risk is that they may make structural injustice easier to absorb.*

Representation is Necessary, but Not Sufficient

Current debates around AI and disability often focus on representation. Researchers and advocates have highlighted how AI systems fail disabled users when disability experiences are poorly represented in training data.

Improving representation in datasets is therefore essential. But representation alone does not resolve the issue described here.

Even when disability narratives are present in training data, AI systems still operate through processes of summarization and pattern detection. This process is inherently compressive. Pattern detection privileges recurrence. Structured outputs favor coherence and balance.

Adding more disability data improves recognition. *It does not guarantee preservation.*

Knowledge Systems and Interpretive Power

Scholars have long recognized that the frameworks through which testimony is interpreted shape how knowledge becomes legible. Philosopher Miranda Fricker describes how individuals can be wronged (<https://academic.oup.com/book/32817>) in their capacity as knowers when prejudice on the part of hearers leads their testimony to be undervalued or dismissed. In such cases, the problem is not factual inaccuracy but a distortion in how credibility and meaning are assigned.

Research in science and technology studies extends this insight to institutional systems of knowledge organization. Geoffrey Bowker and Susan Leigh Star showed how classification systems determine what becomes visible (<https://mitpress.mit.edu/9780262522953/sorting-things-out/>) within knowledge infrastructures and what remains obscured.

AI writing assistants introduce yet another layer into these interpretive processes. While they are not moral agents capable of prejudice, they operate through statistical patterns and training data that can shape how testimony is summarized, categorized, and represented. In doing so, they help shape how lived experience becomes research knowledge.

In such processes, the issue is rarely a factual error but rather *how emphasis, tone, and moral weight are redistributed*.

Guardrails for AI Tools in Scholarly Publishing

None of this suggests that AI writing assistants should be rejected. AI tools can help researchers navigate complex literature, improve clarity, and reduce the mechanical burden of manuscript preparation.

Much of the current debate focuses on risks such as hallucinated sources, biased training data, and potential deskilling of researchers. But the experiment described above highlights a different issue: how AI-mediated summarization may subtly redistribute the emotional and interpretive weight carried by qualitative testimony.

As publishers integrate AI into author workflows, design choices will shape how research knowledge is represented.

A few design guardrails may help ensure that AI-assisted writing tools support rather than inadvertently dilute qualitative scholarship:

- **Preserve participant voice in qualitative research.** Author tools that generate summaries should not replace direct testimony. In many forms of qualitative research, participant quotations carry analytic and ethical significance that cannot be captured through thematic synthesis alone.

- **Avoid automatic normalization of critique.** Systems trained to produce balanced professional prose may unintentionally soften structural critique. Tool design should prioritize clarity without automatically translating emotionally charged testimony into neutral administrative language.
- **Make summarization visible.** AI-generated summaries inevitably compress complex material. Tools should make this transformation visible, so authors understand that summaries represent interpretive reductions rather than neutral representations of evidence.
- **Support interpretive accountability.** AI writing assistants should be designed to support authors rather than replace analytic judgment. Deciding which experiences carry significance and how they should be represented remains a fundamentally human responsibility.

Interpretation Cannot Be Automated

Scholarly publishing has always involved mediation. Authors interpret evidence. Editors shape narratives. Peer reviewers evaluate claims.

AI writing assistants now add a new layer to that process. These tools can help researchers see patterns across large bodies of material. They can improve clarity and reduce the mechanical burden of writing.

But they cannot determine which experiences carry moral significance. They cannot decide which edges of testimony should remain sharp.

As AI becomes part of the infrastructure of scholarly communication, from literature discovery to manuscript drafting, preserving the interpretive integrity of lived experience will remain a responsibility that cannot be automated.

AI can help us see patterns. But deciding what those patterns mean is still human work.

Author's note: For the exploratory exercise described in this post, I used Microsoft Copilot, accessed through a web browser. The purpose was not to evaluate Copilot as a specific product, but to use a widely available generative AI system to explore how AI-mediated summarisation may affect qualitative testimony. The observations in this post should therefore be read as a broader reflection on AI-assisted summarization and research writing, rather than as a product review of Copilot.

Sylwia Frankowska-Takhari

Dr. Sylwia Frankowska-Takhari is an accessibility specialist at Springer Nature, where she supports teams with accessibility audits, inclusive design guidance, and accessibility strategy. She has over 15 years of experience in user research, digital accessibility, and inclusive design across higher education, publishing, finance, and international development. Sylwia is also a visiting lecturer in human-computer interaction and inclusive design at City St George's, UoL, and her work focuses on accessibility, disabled user experience, and the role of research in shaping more inclusive products and services.

Discussion

4 THOUGHTS ON "GUEST POST – WHEN AI HELPS WRITE RESEARCH: WHAT HAPPENS TO LIVED EXPERIENCE?"



It would be possible, of course, in the prompt to the AI system, to tell the system to use the personal language style found in the quotations, rather than abstract reformulations of these words. When I use AI (usually Claude.ai) I give instructions of this kind and find that the style is adjusted.

By TOM WILSON | MAY 21, 2026, 9:30 AM

@ TOM WILSON That was my immediate thought as well. Frankowska-Takhari raises important points, but I don't think that those are particularly difficult to account for via prompting and other tools that change defaults. Microsoft has instructions for Copilot here:

<https://learn.microsoft.com/en-us/microsoft-copilot-studio/guidance/generative-mode-guidance>

By THADDEUS MCILROY | MAY 21, 2026, 12:53 PM

Thank you both! Yes, I agree that prompting and tool configuration can make a significant difference. I don't see the issue as "AI cannot do this", but rather that users may not realise they need to ask for this kind of preservation by default. In accessibility and disability-related contexts especially, we may need to be quite intentional about asking AI systems to preserve specificity, direct language, and the meaning carried by testimony, not only the broad themes.

By SYLWIA FRANKOWSKA-TAKHARI | MAY 21, 2026, 1:50 PM



Yes, probably the majority of AI users are unaware of how the nature of the prompt can influence the output – this is part of what has come to be called AI-literacy, which clearly does not reach the population at large. A further thought: I use Claude.ai and it is possible to prepare a "skill" document which Claude.ai uses in interactions – this sets out the criteria to be applied when preparing outputs covering everything from language, formatting of text, etc., etc. Claude can draft such a document on the basis of past interactions which the user can then edit. The use of a "skill" in this way can ensure that the user receives output in a desired form.

By TOM WILSON | MAY 22, 2026, 6:48 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