

Epistemic Transformations: Large Language Models and the Reconfiguration of Scholarly Knowledge Production

Jonathan H. Westover^{a b c *}

^a Utah Valley University, UT, USA

^b Future State University, CA, USA

^c Nexus Institute for Work & AI, Catalyst Center for Work Innovation, UT, USA

* Correspondence: jon.westover@gmail.com

Received April 15, 2026

Accepted for publication July 17, 2026

Published Early Access August 11, 2026

doi.org/10.70175/innovationjournal.2026.1.1.4

Abstract: *The rapid integration of Large Language Models (LLMs) into academic research practices presents significant opportunities and challenges that extend beyond questions of efficiency to fundamental reconfigurations of scholarly knowledge production itself. This article provides a comprehensive, interdisciplinary examination of how these artificial intelligence systems are reshaping research workflows, epistemic practices, and the very categories through which we understand scholarship. Drawing on Science and Technology Studies, philosophy of science, philosophy of mind, research ethics, and scholarship from diverse global contexts, the analysis situates LLM integration within broader debates about technological mediation in knowledge production while attending to how these technologies may transform the foundational concepts of authorship, originality, and expertise. The article examines empirical evidence regarding adoption patterns, critically assesses both benefits and risks—including concerns about epistemic quality, research integrity, equity, and the preservation of scholarly competencies—while engaging systematically with counterarguments and alternative perspectives. Particular attention is given to disciplinary variation through detailed case studies, global and linguistic dimensions drawing on non-Anglophone scholarship, and temporal dynamics of adoption. A multi-level framework for responsible integration is proposed, offering operationalized guidance for individual researchers, institutions, publishers, and policymakers, alongside examination of coordinated governance mechanisms. The framework addresses the fundamental tension between leveraging technological capabilities and maintaining the epistemic virtues that underpin trustworthy scholarship, while acknowledging that these very virtues may themselves be undergoing transformation. This analysis contributes to ongoing debates about the future of academic knowledge production by providing theoretically grounded, empirically informed, and practically oriented guidance for navigating this transformative technological moment.*

Keywords: artificial intelligence, large language models, scholarly communication, research ethics, academic writing, knowledge production, research integrity, epistemic cultures, epistemic justice, technological mediation

Suggested Citation:

Westover, Jonathan H. (2026). Epistemic Transformations: Large Language Models and the Reconfiguration of Scholarly Knowledge Production. *Innovation: The Journal of Knowledge Work, Creativity, and Organizational Advancement*, 1(1). doi.org/10.70175/innovationjournal.2026.1.1.4

1. Introduction

The emergence of Large Language Models (LLMs) represents a technological development whose implications for scholarship may extend beyond enhanced efficiency to fundamental transformations in how knowledge is produced, validated, and understood. Since the public release of ChatGPT in November 2022, researchers across disciplines have begun experimenting with these systems for tasks ranging from literature review and data analysis to manuscript drafting and code generation (Stokel-Walker & Van Noorden, 2023). This rapid adoption has prompted urgent discussions among scholars, institutions, publishers, and policymakers about appropriate use, disclosure requirements, and the preservation of research integrity.

Yet framing LLM integration solely in terms of "appropriate use" may itself reflect assumptions that warrant examination. Such framing presupposes stable categories—author, tool, original contribution—into which LLMs must be fitted. A deeper analysis must consider whether LLMs are not merely new tools to be governed but technologies that may reconfigure the very categories through which scholarship has traditionally been understood. This article takes seriously both practical questions of governance and the more fundamental question of whether our existing conceptual apparatus remains adequate.

Understanding the implications of LLM integration constitutes an important scholarly endeavor for several interconnected reasons. First, the pace of adoption has outstripped the development of normative frameworks, creating uncertainty about appropriate practices and potential for inconsistent standards across disciplines and institutions. Second, the capabilities of these systems continue to evolve rapidly, requiring ongoing reassessment of both opportunities and risks. Third, decisions made during this formative period may establish path dependencies that shape scholarly practices for decades. Fourth, the global nature of academic research means that adoption patterns and their consequences will affect knowledge production worldwide, with potentially differential impacts across regions and scholarly traditions. Fifth, and perhaps most fundamentally, the integration of LLMs raises questions about what scholarship *is*—questions that have been relatively settled but may now require reopening.

This article provides a comprehensive, interdisciplinary examination of LLM integration in scholarly production. The analysis proceeds in several stages. Following this introduction, Section 2 establishes conceptual and theoretical foundations, engaging deeply with Science and Technology Studies frameworks not merely as vocabulary but as analytical resources that challenge conventional framings. Section 3 examines methodological approaches and acknowledges limitations. Section 4 surveys LLM applications across the research lifecycle, with detailed disciplinary case studies. Section 5 critically analyzes implications for epistemic quality and research integrity. Section 6 addresses equity considerations with particular attention to global dimensions, drawing on non-Anglophone scholarship. Section 7 examines emerging policy responses and institutional frameworks, including coordinated governance mechanisms. Section 8 proposes an operationalized framework for responsible integration. Section 9 discusses temporal dynamics and future trajectories, and Section 10 offers concluding reflections.

The central argument advanced is that LLMs function as more than cognitive tools that augment human research capabilities—they are technologies that participate in and potentially transform the networks through which scholarly knowledge is produced. This transformation may ultimately reconfigure foundational concepts including authorship, originality, and expertise. Rather than offering definitive answers to questions that remain genuinely open, this article aims to provide conceptual resources and practical frameworks for navigating these complexities while remaining attentive to the possibility that the very frameworks we construct may require revision as our understanding deepens.

2. Conceptual and Theoretical Foundations

2.1 Defining Large Language Models

Large Language Models are artificial intelligence systems trained on extensive text corpora to predict and generate human-like text (Brown et al., 2020). These models employ transformer architectures that process sequential data through attention mechanisms, enabling them to capture complex linguistic patterns and relationships across extended contexts (Vaswani et al., 2017). Contemporary LLMs such as GPT-4, Claude, and Gemini contain billions of parameters and have demonstrated capabilities in text generation, summarization, translation, and various forms of reasoning-like behavior.

Several technical characteristics bear particular relevance to scholarly applications. First, these systems exhibit *emergent capabilities*—abilities not explicitly programmed but arising from scale and training—including apparent competence in specialized domains, code generation, and multi-step problem solving (Wei et al., 2022). Second, LLMs demonstrate *contextual adaptability*, adjusting outputs based on prompts, instructions, and conversational context. Third, their *knowledge integration* spans diverse domains, reflecting training on heterogeneous text sources including scientific literature, though with significant temporal, linguistic, and geographic limitations.

2.2 Science and Technology Studies: Beyond Vocabulary to Analytical Transformation

Science and Technology Studies (STS) provides essential analytical resources for understanding how LLMs may reshape scholarly practice—but engaging seriously with STS requires more than borrowing its terminology. STS scholarship challenges the very distinctions that frame conventional discussions of technology, including the boundaries between human and nonhuman agency, between tool and user, and between the social and the technical (Latour, 2005; Pickering, 1995). Applied to LLMs in scholarship, this means questioning whether our existing categories—author, instrument, original contribution—remain adequate or whether LLM integration reveals their contingency.

2.2.1 Technological Mediation and the Transformation of Practice

The concept of *technological mediation* proves particularly illuminating when developed beyond its surface implications. Ihde's (1990) postphenomenology describes how technologies mediate human-world relations through different structures: *embodiment relations* (where technology becomes transparent, like eyeglasses), *hermeneutic relations* (where we read the world through technology, like a thermometer), and *alterity relations* (where technology appears as a quasi-other with which we interact).

LLMs seem to occupy an unstable position across these categories. In some uses—grammar checking, for instance—they may approach the transparency of embodiment relations. In others—conversational interaction for brainstorming—they function more like quasi-others in alterity relations. In still others—generating text that a researcher then edits—the relationship may be something new that strains existing categories. This categorical instability itself is analytically significant: LLMs may not fit neatly into our existing frameworks for understanding human-technology relations.

Crucially, technological mediation is not neutral. Technologies do not merely facilitate pre-existing intentions but shape what becomes thinkable, sayable, and doable. If LLMs mediate scholarly writing, they participate in shaping not just how scholars express ideas but what ideas get developed. This is not necessarily problematic—all writing technologies have such effects—but it warrants attention, particularly given the opacity of LLM processes and the homogenizing tendencies that may emerge from training on similar corpora.

2.2.2 Actor-Network Theory and the Redistribution of Agency

Actor-Network Theory (ANT) offers additional analytical resources by treating human and non-human entities symmetrically within networks of action (Latour, 2005; Callon, 1986). From this perspective, LLMs become *actants*—entities that make differences in networks of knowledge production. The principle of *generalized symmetry* does not claim that humans and LLMs are identical but rather that the analyst should not assume *a priori* that agency, intentionality, or cognitive capacity are exclusively human properties that explain outcomes.

Applying ANT to LLM-assisted scholarship means tracing how the introduction of LLMs redistributes agency within research networks. Questions shift from "Is this appropriate LLM use?" to "How does LLM involvement reconfigure the network of relations through which this knowledge claim is produced, validated, and circulated?" This reframing has practical implications:

- **Authorship** becomes a question about network configuration rather than individual contribution
- **Verification** becomes distributed across human-LLM assemblages rather than located in individual cognition
- **Responsibility** must be renegotiated as agency becomes less clearly locatable

This perspective does not dissolve normative questions but reframes them. The question is not whether to "allow" LLMs in some pristine human practice but how to configure human-LLM networks in ways that produce trustworthy knowledge while maintaining meaningful accountability.

2.2.3 Boundary Objects and Epistemic Translation

Star and Griesemer's (1989) concept of *boundary objects* illuminates additional LLM functions in scholarly work. Boundary objects are entities that inhabit multiple social worlds and satisfy the informational requirements of each, maintaining coherence across contexts while being adaptable within them. LLMs may function as boundary objects by facilitating communication across disciplinary boundaries, translating specialized concepts, and enabling collaboration between researchers with different expertise.

However, this boundary-crossing capacity carries risks that STS analysis helps identify. When boundary objects enable translation between communities, they may also flatten important distinctions. If an LLM translates between, say, humanistic and quantitative approaches to a problem, it may do so in ways that obscure the incommensurabilities that make each approach distinctive. Disciplines are not merely different vocabularies for the same underlying reality; they embody different epistemologies, different relationships to evidence, different standards of demonstration. LLMs that facilitate interdisciplinary work may inadvertently impose false commensurability.

2.2.4 Reconfiguring Foundational Categories

A thoroughgoing STS analysis must consider whether LLMs call into question the foundational categories through which scholarship has been understood:

Authorship: The Romantic conception of authorship as individual creative genius has long been criticized (Barthes, 1967; Foucault, 1984), but scholarship has retained attribution practices that assume identifiable individual contributions. LLMs challenge this by introducing a form of contribution that is neither individual (LLMs synthesize training from countless sources) nor collaborative in the traditional sense (LLMs do not negotiate meaning). We may need new conceptual resources—perhaps something like "assembled authorship" or "distributed composition"—to describe what LLM-assisted scholarship involves.

Originality: Scholarly contribution traditionally requires original thought—novel arguments, new findings, creative interpretations. But if LLMs generate text by recombining patterns from training data, what happens to originality when human authors work extensively with such text? The question is not simply whether LLM outputs are original (they are arguably not, in the traditional sense) but whether the human work of prompting, selecting, editing, and taking responsibility constitutes a different form of originality that remains valuable.

Expertise: Expertise has traditionally involved both propositional knowledge and tacit understanding developed through experience. LLMs possess something like propositional knowledge (they can articulate information accurately or inaccurately) but not tacit understanding grounded in practice. Their involvement in scholarly work may reconfigure what counts as expertise—perhaps shifting emphasis from information possession (which LLMs can assist with) to judgment, interpretation, and accountability (which they cannot).

This analysis suggests that responsible LLM integration is not merely a matter of fitting new tools into existing categories but of reconceptualizing categories themselves. This is unsettling but also potentially productive: periods of technological change often occasion valuable reflection on previously tacit assumptions.

2.3 Philosophical Considerations: Understanding, Intentionality, and Epistemic Humility

A critical question for evaluating LLM integration concerns the nature and limits of their cognitive capacities. Early versions of this analysis asserted that LLMs lack "genuine understanding," but this claim requires careful philosophical examination.

2.3.1 The Chinese Room and Its Discontents

The most influential philosophical challenge to machine understanding comes from Searle's (1980) Chinese Room argument. Searle argued that a system manipulating symbols according to formal rules—as LLMs do—could produce outputs indistinguishable from those of an understanding agent without possessing genuine understanding itself. The crucial distinction is between *syntax* (manipulation of symbols according to rules) and *semantics* (meaning and reference). For Searle, computers operate purely syntactically, while understanding requires semantics.

However, the Chinese Room argument has generated extensive critical responses (Cole, 2020). The *Systems Reply* argues that while no component of the room understands, the system as a whole might. The *Robot Reply* suggests that understanding requires embodiment and environmental interaction, which the thought experiment artificially excludes. The *Brain Simulator Reply* questions whether the intuition would hold if the room simulated neurons rather than conversation.

More fundamentally, philosophers of cognitive science have questioned whether "understanding" names a single, well-defined capacity or a family of loosely related abilities (Dennett, 1991; Clark, 2008). If understanding comes in degrees and varieties, the question "Do LLMs understand?" may be malformed. More productive questions might concern which cognitive capacities LLMs possess, to what degree, and what implications these capacities have for their scholarly applications.

2.3.2 Functional Considerations

Regardless of metaphysical questions about machine minds, several functional observations are relevant for scholarly contexts:

1. **Epistemic uncertainty:** We cannot currently determine with confidence whether LLMs possess anything like understanding. The philosophical question remains genuinely open, and intellectual honesty requires acknowledging this uncertainty.
2. **Functional limitations:** LLMs currently exhibit limitations that matter for scholarship—including unreliable factual accuracy, inability to verify their outputs against the world, lack of access to non-textual evidence, and absence of the lived experience that grounds much humanistic and social scientific interpretation.
3. **Accountability asymmetry:** Whatever their cognitive status, LLMs cannot be held responsible for scholarly claims in the way humans can. They cannot explain their reasoning in ways that reveal genuine justification, cannot respond to criticism by reconsidering their views, and cannot be sanctioned for misconduct. This accountability asymmetry has practical implications regardless of metaphysical questions.
4. **Precautionary stance:** Given uncertainty about LLM capacities, prudence suggests treating their outputs as requiring human verification and judgment. If LLMs turn out to be more capable than we assume, little is lost by excessive caution; if they are less capable, excessive trust could cause significant harm.

This framework acknowledges philosophical humility while still providing grounds for practical recommendations. The recommendations do not depend on resolving metaphysical questions but on recognizing functional limitations and accountability structures.

2.4 Empirical Evidence on Adoption Patterns: What We Know and Don't Know

Understanding actual adoption patterns provides essential context for assessing implications. However, intellectual honesty requires acknowledging that the empirical evidence base remains limited.

2.4.1 Available Evidence

Survey research has begun documenting adoption patterns. Gao et al. (2022) analyzed language patterns in scientific abstracts and estimated increased use of AI writing assistance based on stylistic markers, though such indirect methods have interpretive limitations. The *Nature* survey of researchers found widespread experimentation with LLMs for various tasks, with significant variation by career stage, discipline, and geography (Van Noorden & Perkel, 2023). Early-career researchers reported higher adoption rates, consistent with broader patterns in technology diffusion. Lund and Wang (2023) documented diverse practices in academic contexts, ranging from minimal editing assistance to more substantial text generation.

2.4.2 Evidentiary Limitations

Much of the evidence regarding adoption patterns and effects comes from:

- Self-report surveys with potential selection bias (researchers willing to respond may differ systematically from non-respondents)
- Indirect linguistic analyses with contested interpretation
- Anecdotal reports and journalistic accounts
- Early preprints not yet subjected to peer review
- English-language sources that may not represent global patterns

This evidentiary landscape suggests caution in making strong empirical claims. Throughout this article, I distinguish between claims supported by robust evidence, those based on emerging but limited evidence, and those that represent theoretical arguments or informed speculation.

2.5 Epistemic Cultures: Disciplinary Variation as More Than Surface Difference

Knorr-Cetina's (1999) concept of *epistemic cultures* provides crucial resources for understanding how LLM integration varies across disciplines. Different fields possess distinct "machineries of knowledge construction"—different relationships to instruments, different forms of evidence, different authorship practices, and different standards of demonstration. These differences are not merely methodological preferences but reflect deep-seated epistemological commitments about what knowledge is and how it is properly produced.

Taking epistemic cultures seriously means recognizing that LLM integration cannot be addressed through uniform policies. What counts as appropriate use depends on what a discipline values and how it produces knowledge. Later sections develop detailed case studies (Section 4.6) that illustrate this variation concretely.

3. Methodology and Limitations

3.1 Analytical Approach

This article represents a critical narrative synthesis drawing on multiple literatures: Science and Technology Studies, philosophy of science and mind, research ethics, information science, and the emerging empirical literature on AI in research contexts. Literature was identified through systematic database searches (Web of Science, Scopus, Google Scholar), tracking citations in key sources, monitoring relevant preprint servers (arXiv, OSF Preprints), and following professional discussions in relevant venues.

Recognizing the limitations of English-language scholarship noted by reviewers, targeted efforts were made to identify relevant non-Anglophone scholarship, including Spanish-language work on AI in Latin American academic contexts (Fischman & Alperin, 2015, on Latin American scholarly communication more broadly), French-language contributions to technology studies, and scholarship addressing Asian academic contexts (Flowerdew, 2019, on linguistic disadvantage). These efforts remain incomplete—a genuine limitation acknowledged below—but represent an attempt to broaden the evidentiary base.

The analytical approach is primarily theoretical, applying established frameworks to the novel phenomenon of LLM integration. Empirical claims are incorporated where available but distinguished from theoretical arguments. The article also incorporates the kind of reflexivity that STS scholarship commends: the analysis itself is shaped by the author's position, disciplinary training, linguistic capacities, and context.

3.2 Limitations

Several limitations should be acknowledged:

1. **Rapidly evolving landscape:** LLM capabilities and adoption patterns change faster than traditional publication cycles. Some observations may be outdated by publication.
2. **Limited empirical base:** Much discussion of LLM effects remains speculative. Robust empirical research on actual impacts is only beginning to emerge.
3. **English-language predominance:** Despite efforts to incorporate non-Anglophone scholarship, the analysis draws primarily on English-language sources. This reflects both the author's linguistic

limitations and the structure of academic databases, but it means perspectives from other scholarly traditions are underrepresented.

4. **Disciplinary limitations:** While the article addresses disciplinary variation through case studies, the author's expertise does not span all fields equally. The case studies represent informed interpretations that scholars from those disciplines might contest or refine.
5. **Focus on text-based LLMs:** The analysis centers on text-generating models; multimodal systems and domain-specific AI raise considerations not fully addressed here.
6. **Normative uncertainty:** The article proposes frameworks for "responsible" integration, but what counts as responsible remains contested. The frameworks offered are contributions to an ongoing conversation rather than final determinations.

4. LLM Applications Across the Research Lifecycle

4.1 Literature Review and Synthesis

Literature review represents a foundational research activity where LLMs offer potentially significant assistance. Traditional comprehensive literature review requires substantial time investment and expertise in locating, evaluating, and synthesizing sources. LLMs can accelerate aspects of this process by summarizing articles, identifying themes across large corpora, and suggesting relevant sources.

Potential benefits include enabling broader coverage than time constraints typically allow, assisting researchers venturing into unfamiliar literatures, and facilitating identification of patterns across large bodies of work.

Concerns center on whether LLM-assisted review achieves the same epistemic goals as traditional practices. When researchers engage deeply with primary sources, they develop tacit knowledge about a field—understanding not just what has been claimed but how claims are supported, contested, and contextualized within ongoing scholarly conversations. LLM summaries, even when accurate, may not support this form of learning.

Illustrative example: A researcher uses an LLM to summarize 50 articles on sustainable urban development. The LLM produces accurate summaries of each paper's claims. But the researcher may miss the methodological disagreements between quantitative and qualitative researchers in this space, the tensions between engineering and social science approaches, and the ways that "sustainability" itself is contested and differently operationalized across studies. These absences matter: they constitute the kind of knowledge that enables sophisticated contribution to a scholarly conversation rather than mere citation of prior work.

Disciplinary variation: In humanities disciplines where engagement with primary texts is epistemically central—literary studies, historical research, philosophical analysis—LLM-mediated literature review may be more problematic than in fields where the goal is primarily to establish the current state of knowledge on specific empirical questions.

4.2 Hypothesis Generation and Research Design

LLMs can assist in brainstorming research questions, identifying gaps in existing literature, and suggesting methodological approaches. By processing extensive literature, these systems may identify connections or unexplored territory that researchers might overlook.

However, the quality of LLM-generated research directions remains uncertain. Systematic evaluation of whether LLM suggestions lead to productive research is lacking.

Counterargument considered: One might argue that researchers need not generate all ideas themselves—insights can come from diverse sources, including conversations with colleagues, serendipitous discoveries, and now AI systems. The source of an idea matters less than its subsequent development and testing. Scientists have long used various prompts and assistants; LLMs are simply another resource.

Response: This argument has merit but may understate important differences. Ideas that emerge from deep engagement with problems, materials, and disciplinary traditions are shaped by that engagement in ways that may make them more likely to be productive and to connect with ongoing scholarly conversations. More importantly, the capacity for hypothesis generation is part of what we mean by expertise; if this capacity atrophies through delegation to LLMs, something significant may be lost even if individual ideas prove valuable.

The honest assessment is that we do not yet know whether LLM-assisted hypothesis generation leads to better, worse, or different research. Longitudinal studies would be valuable.

4.3 Data Analysis Support

In quantitative research, LLMs increasingly assist with writing analysis code, explaining statistical outputs, and suggesting analytical approaches. Researchers can describe their data and analytical goals in natural language and receive code in R, Python, or other languages (Chen et al., 2021).

Benefits: This capability can democratize advanced methods, enabling researchers without extensive programming training to implement sophisticated analyses. It may accelerate workflows for experienced analysts and reduce the barrier between conceptualizing and implementing analyses.

Risks: Generated code may contain errors that are not immediately apparent to users unfamiliar with the underlying methods. Researchers may apply techniques without fully understanding their assumptions and limitations. The efficiency gains might come at the cost of methodological competence.

Case example: A health policy researcher with limited statistical training receives reviewer requests for multilevel modeling. Using an LLM, she generates code that runs without errors and produces plausible output. She includes the analysis in her revision. However, she cannot fully evaluate whether the model specification is appropriate, whether assumptions are met, or how to interpret the variance components. The paper is published; readers assume the analysis is sound because it appears in a peer-reviewed venue.

Analysis: This scenario illustrates a gap between technical capability (running the code) and genuine competence (understanding the analysis). Disclosure of LLM assistance does not resolve the underlying issue. Possible mitigations include collaboration with methodological experts, explicit acknowledgment of analytical limitations, or reviewer capacity to evaluate LLM-assisted analyses.

4.4 Writing and Manuscript Preparation

Manuscript drafting represents the most visible and contentious LLM application. Uses range from sentence-level editing to generating substantial text portions. Practices exist along a spectrum:

Minimal use: Grammar and spelling correction, minor stylistic improvements

Moderate use: Paragraph-level revision, restructuring, translating notes into prose

Substantial use: Generating first drafts of sections based on detailed prompts

Maximal use: Generating entire manuscripts with minimal human input beyond prompting

Distinguishing legitimate from problematic use is challenging because the same tools can serve different purposes, and the distinction lies not in the technology but in the human engagement around it. A proposed criterion: use is appropriate when the human author genuinely understands and can defend all claims made, can explain the reasoning behind arguments, and takes full intellectual responsibility for the work. Use is problematic when the human author could not articulate or defend the claims without reference to the LLM that generated them.

This criterion is difficult to operationalize and impossible to verify externally, which is precisely why detection-focused enforcement is likely to be limited. The criterion is better suited to guiding individual conduct and professional norms than to external policing.

4.5 Peer Review

LLM use in peer review remains controversial. Potential applications include assisting reviewers with literature checks, identifying methodological issues, or structuring reviews. Checco et al. (2021) documented early explorations, though these preceded current LLM capabilities.

The peer review system depends on expert human judgment regarding validity, significance, and contribution. While LLMs might assist with routine aspects (checking reference formatting, identifying potential plagiarism), core evaluative functions seem to require human expertise.

Emerging concerns: Reports suggest some reviewers are using LLMs to generate or substantially draft reviews. This raises concerns about review quality (LLMs may miss subtle issues or generate generic comments), confidentiality (submitting manuscripts to external LLM services may expose unpublished work), and the integrity of the peer review process itself.

Publishers have begun implementing policies restricting or guiding LLM use in review, though enforcement mechanisms remain unclear.

4.6 Disciplinary Case Studies

To move beyond general observations, this section examines LLM integration in three disciplines in depth, illustrating how epistemic cultures shape appropriate use.

4.6.1 Case Study: History

Historical scholarship involves distinctive epistemic practices: engagement with primary sources (archival documents, material artifacts, visual evidence), interpretation of sources within their historical contexts, construction of narratives that explain change over time, and reflexive attention to the historian's own positionality and interpretive frameworks.

How might LLMs be used?

- Transcription assistance for handwritten documents
- Translation of sources in languages the historian does not read fluently
- Summarizing large volumes of secondary literature
- Drafting routine prose (e.g., contextual background)

What concerns arise?

First, *interpretive depth*: Historical understanding emerges from close engagement with sources—noticing inconsistencies, reading against the grain, attending to absences. LLM summaries cannot replicate this engagement and may short-circuit the interpretive process through which historical knowledge is produced.

Second, *voice and narrative*: Historical writing often involves distinctive authorial voice and narrative construction. LLM-generated prose may homogenize historical writing, losing the diversity of styles and perspectives that characterize the field.

Third, *source verification*: LLMs may fabricate sources or misrepresent their content. Given that historical claims depend on evidentiary grounding, this is particularly concerning. A historian who cites a source based on an LLM summary without consulting the original risks serious error.

Fourth, *tacit knowledge*: Experienced historians develop tacit knowledge about their periods—familiarity with names, events, social contexts—that informs interpretation. This knowledge develops through sustained engagement with sources. If LLMs substitute for such engagement, this knowledge may not develop.

Implications for the framework: In historical scholarship, LLM use for source engagement should be approached cautiously. Use for peripheral tasks (formatting, grammar) is less concerning than use for interpretive work central to historical contribution. Verification requirements are particularly stringent given the evidentiary nature of historical claims.

4.6.2 Case Study: Biomedical Research

Biomedical research combines laboratory practices, clinical observations, and increasingly computational and data-intensive approaches. Epistemic cultures within biomedicine are themselves diverse—molecular biology differs from epidemiology, which differs from clinical trials research.

How might LLMs be used?

- Literature review across a vast and rapidly growing corpus
- Analysis code for computational biology and bioinformatics
- Drafting methods sections with standardized protocols
- Assisting non-native English speakers in writing for international journals

What concerns arise?

First, *patient safety implications*: Biomedical research often has downstream implications for patient care. Errors introduced through LLM assistance could propagate to clinical recommendations. The stakes of inaccuracy are particularly high.

Second, *data integrity*: LLMs trained on published literature may perpetuate errors, biases, or superseded findings present in that literature. The replication crisis in biomedical research suggests that substantial published literature contains errors; LLM integration should not amplify this problem.

Third, *authorship practices*: Biomedical research often involves large collaborative teams with established norms about contribution and authorship. LLM involvement complicates these already complex attribution questions.

Fourth, *regulatory considerations*: Research involving human subjects, drug development, or medical devices operates within regulatory frameworks. Regulators are beginning to address AI involvement; researchers must attend to emerging requirements.

Implications for the framework: In biomedical research, verification requirements are particularly stringent given patient safety implications. Transparency in disclosure is important for regulatory and ethical reasons. Collaborative verification—having multiple team members check LLM-assisted work—may be appropriate. Particular caution is warranted for claims that might directly affect clinical practice.

4.6.3 Case Study: Literary Studies

Literary scholarship centers on interpretation of texts, attention to language and form, and engagement with theoretical frameworks that shape how texts are read. The field values original interpretation, close reading, and distinctive critical voice.

How might LLMs be used?

- Locating passages in long texts relevant to specific themes
- Identifying patterns across large corpora (computational literary studies)
- Summarizing secondary criticism
- Editing prose for clarity

What concerns arise?

First, *interpretation vs. pattern recognition*: Literary interpretation involves more than identifying patterns; it involves arguing for the significance of patterns, connecting them to broader concerns, and offering readings that illuminate texts in new ways. LLMs can identify patterns but not argue for significance in ways grounded in aesthetic judgment and theoretical commitment.

Second, *critical voice*: Literary criticism often involves distinctive authorial voice as part of the scholarly contribution. LLM-generated prose may be competent but generic, lacking the voice that characterizes significant criticism.

Third, *theoretical engagement*: Literary scholars work within and against theoretical traditions—postcolonialism, feminism, deconstruction, etc. Genuine engagement with these traditions involves understanding their histories, debates, and commitments. LLMs may reproduce theoretical vocabulary without genuine understanding, producing what appears to be theoretically informed criticism but lacks depth.

Fourth, *computational literary studies*: The subfield of computational literary studies has already normalized AI and computational methods for textual analysis. LLM integration may be less disruptive here, though questions about interpretation vs. pattern recognition remain relevant.

Implications for the framework: In literary studies, LLM use for interpretive claims central to contributions is particularly problematic. Use for peripheral tasks is less concerning. Scholars should be cautious about LLM-generated theoretical claims, which may be superficially plausible but lack genuine engagement with theoretical traditions. In computational literary studies, LLMs may be more readily integrated but should complement rather than substitute for interpretive judgment.

5. Implications for Epistemic Quality and Research Integrity

5.1 Factual Reliability and the Hallucination Problem

LLMs can generate plausible but inaccurate content, a phenomenon termed "hallucination" (Ji et al., 2023). For scholarship, this poses significant risks: fabricated citations, incorrect statistical claims, or unfounded assertions might enter the literature if researchers do not carefully verify LLM outputs.

The hallucination problem is particularly insidious because LLM outputs are presented with uniform confidence regardless of accuracy. Unlike human experts who typically hedge uncertain claims or acknowledge limitations, LLMs may assert fabrications with the same fluency as accurate information.

Counterargument considered: Citation errors, factual inaccuracies, and even fabrication occur in scholarship without LLM involvement. Studies have documented substantial rates of citation errors in published literature (Simkin & Roychowdhury, 2003). LLMs may not be qualitatively worse—just a different source of error.

Response: This argument has partial merit but may understate important differences. Human errors typically occur despite scholars' intentions and efforts to be accurate; they result from oversight, memory failure, or misunderstanding. LLM hallucinations are a structural feature of how these systems work—they generate plausible text rather than retrieve verified facts. The mechanisms differ even if the surface manifestations are similar.

Additionally, the scale at which LLMs can produce content means that the aggregate volume of potential errors could be much greater than under traditional practices. A human might produce one paper with an erroneous citation; an LLM might assist with hundreds of papers, each potentially containing fabricated references.

Finally, LLM integration shifts the verification burden. Previously, authors were responsible for accuracy in claims they generated. Now authors must verify claims generated by a system whose processes they do not fully understand. This verification may be less reliable than original accuracy.

5.2 Bias Propagation in Scholarly Contexts

LLMs encode biases from training data, potentially perpetuating problematic assumptions. Ferrara (2023) examined bias in AI systems and highlighted risks of perpetuating harmful stereotypes or perspectives. In scholarly contexts, bias propagation might occur through several mechanisms:

- **Citation bias:** LLMs may over-represent frequently-cited sources, reinforcing existing hierarchies and under-representing emerging or marginalized scholarship (Thelwall, 2022).
- **Theoretical bias:** Dominant theoretical frameworks in training data may be reproduced, potentially marginalizing alternative approaches.
- **Methodological bias:** Training data may skew toward certain methodological traditions, affecting how research questions are framed or what approaches are suggested.
- **Disciplinary bias:** Fields better represented in training data may receive more sophisticated assistance than underrepresented disciplines.
- **Geographic and linguistic bias:** Scholarship from certain regions and in certain languages is underrepresented in training data, leading to biased outputs.

These biases might be less visible than in consumer applications but equally consequential for the direction and inclusivity of knowledge production.

5.3 Implications for Skill Development

Concerns about skill atrophy parallel debates about calculator use in mathematics education. The question is whether offloading tasks to LLMs diminishes capacities that researchers need to develop.

Areas of concern:

- Writing skills developed through struggle with expression
- Critical reading abilities cultivated through direct engagement with literature
- Methodological understanding developed through implementing analyses manually
- Interpretive judgment refined through repeated practice

Counterargument considered: Offloading routine tasks enables focus on higher-order activities. Calculators did not destroy mathematical ability; they enabled engagement with more complex problems. Similarly, LLMs might free researchers for deeper conceptual work. Each technological transition in scholarship—from manuscript to print, typewriter to word processor—prompted similar concerns that proved exaggerated.

Response: This counterargument deserves serious consideration. The historical pattern of technological anxiety followed by successful adaptation provides genuine reassurance. However, several considerations temper this optimism:

First, *the timing of assistance matters*. Using calculators after one has developed arithmetic understanding differs from using them before such understanding develops. For trainees, LLM assistance before developing core skills may be more problematic than assistance afterward.

Second, *the nature of the tasks matters*. Arithmetic is largely mechanical; calculators automate mechanical processes. But writing and interpretation are not merely mechanical—they involve judgment, creativity, and understanding that may require practice to develop. The analogy may be limited.

Third, *the observability of effects matters*. We have decades of evidence about calculator effects on mathematical competence. We have almost no evidence about LLM effects on scholarly competencies. The honest answer is that we do not know whether concerns about skill atrophy will prove warranted. This uncertainty itself suggests caution.

Longitudinal studies comparing researchers who do and do not use LLMs extensively would be valuable. In the meantime, prudent approaches might emphasize developing skills before augmenting them and maintaining periodic unassisted practice.

5.4 Verification Burden and Quality Assurance

If LLMs become widespread in research, the burden of verification shifts significantly. Currently, researchers bear primary responsibility for the accuracy of their outputs. With LLM assistance, researchers must not only produce work but verify LLM contributions that they may not fully understand or could not have generated independently.

This verification burden raises important questions:

- Are the efficiency gains from LLM use partially offset by verification requirements?
- Can researchers reliably identify errors in LLM outputs, especially in areas adjacent to but outside their core expertise?
- What institutional structures might support effective verification?

The effectiveness of verification deserves empirical study. If researchers miss substantial proportions of LLM errors, the claimed benefits of LLM assistance may be illusory—efficiency gains in production offset by errors that escape detection.

5.5 Authorship and Accountability in Reconfigured Networks

Major publishers have concluded that LLMs cannot be listed as authors because they cannot take responsibility for scholarly claims. This position has practical merit: accountability structures require entities that can be held responsible—that can respond to criticism, correct errors, and be sanctioned for misconduct.

However, drawing on the STS analysis developed earlier, the authorship question goes deeper than practical accountability. Foucault's (1984) influential analysis argued that authorship is not simply about who produced a text but about the *author function*—a complex cultural construction involving attribution of meaning, responsibility, and authority. The author function organizes interpretation and establishes relationships that exceed physical text production.

When substantial portions of text are LLM-generated, the author function is destabilized. Several configurations might emerge:

1. **Editorial authorship:** The human provides direction, evaluates outputs, and takes responsibility, but generates little original text. Is this authorship? It resembles traditional authorship less than traditional editing but involves more conceptual work than mere editing implies.
2. **Assembled authorship:** Human and LLM contributions are intermingled in ways difficult to separate. The final text emerges from interaction rather than either party alone.
3. **Curated authorship:** The human selects among LLM-generated options, shapes through prompting, and takes responsibility. Authorship lies in curatorial judgment rather than generation.

None of these fit neatly into traditional authorship concepts. Rather than forcing them into existing categories, we might need new conceptual resources. The key normative criterion might be *substantive responsibility*: Can the author explain and defend all claims? Can they respond to criticism with genuine understanding? Do they endorse the arguments as their own, having genuinely engaged with them rather than merely accepted LLM outputs? This criterion is difficult to operationalize but points toward what matters about authorship beyond mere text production.

6. Equity Considerations: Global Dimensions and Epistemic Justice

6.1 The Access Divide

Unequal access to advanced LLMs threatens to create new stratifications in research capability. Well-resourced institutions can provide researchers access to premium AI services and the computational infrastructure to use them effectively. Less-resourced institutions—disproportionately located in the Global South—may fall further behind.

This extends existing inequalities in research infrastructure. Scholars already disadvantaged by limited library access, computing resources, research funding, or administrative support may find that LLM integration amplifies rather than reduces disadvantage. If LLM assistance becomes normalized in high-resource contexts, scholarship produced without such assistance may be evaluated unfavorably—not because of lesser quality but because of lesser polish.

6.2 Linguistic Dimensions and Non-Anglophone Scholarship

LLM training data skews heavily toward English-language sources. This creates several concerns for global scholarship that extend beyond technical performance.

Performance disparities: LLMs generally perform less well for non-English languages, particularly those with smaller digital footprints. Scholars working in Swahili, Bengali, or Quechua receive less effective assistance than those working in English, German, or French. Even among well-resourced languages, performance varies.

English hegemony reinforcement: LLMs may encourage more researchers to write in English to access better assistance and broader audiences. This accelerates patterns that many scholars have criticized: the marginalization of non-English scholarly traditions, the loss of linguistic diversity in research, and the implicit message that knowledge produced in other languages is less valuable.

Canagarajah (2002) and Flowerdew (2019) have documented how English dominance in academic publishing creates systematic disadvantages for scholars from non-Anglophone backgrounds. LLM integration may intensify these patterns by making English writing even more accessible while leaving other languages behind.

Translation considerations: LLMs can assist with translation, potentially helping non-native English speakers participate in English-dominated venues. This apparent benefit is ambivalent: it may enable individual success while further entrenching systemic English dominance. The "solution" of translating into English does not address the underlying injustice of a system that devalues non-English scholarship.

Loss of linguistic and stylistic diversity: Academic writing in English is already noted for stylistic homogeneity (Hyland, 2004). If LLMs push more scholars toward English and toward LLM-influenced styles, the diversity of scholarly expression may diminish further. Scholarship has aesthetic and cultural dimensions that matter beyond mere information transmission; homogenization represents a genuine loss.

6.3 Knowledge Traditions and Epistemic Justice

Beyond linguistic concerns, there are deeper questions about whose knowledge traditions LLMs represent. Fricker's (2007) work on *epistemic injustice* provides conceptual resources. Epistemic injustice occurs when individuals or groups are wronged in their capacity as knowers—when their testimony is unfairly discredited, or when their experiences lack the conceptual resources for articulation.

If LLM training data underrepresents scholarship from the Global South, Indigenous knowledge traditions, or non-Western intellectual frameworks, LLM assistance may systematically privilege certain epistemologies. Suggestions, summaries, and generated text will reflect dominant knowledge traditions, making alternatives less visible.

Scholars working within underrepresented traditions may find that LLM assistance pushes them toward dominant framings. The very concepts available in LLM outputs may not include those central to their

intellectual projects. This is not merely a matter of incomplete coverage but of potential epistemic violence—the erasure or distortion of knowledge traditions.

Scholars in Latin American contexts have examined how global knowledge hierarchies affect regional scholarship (Fischman & Alperin, 2015; Beigel, 2014). Similar analyses in African, Asian, and other contexts highlight the structures that marginalize non-Western knowledge production. LLM integration should be assessed in light of these concerns, with attention to whether it perpetuates or challenges existing hierarchies.

6.4 Institutional Capacity and Global Governance

Institutions in lower-income countries may lack resources to develop AI policies, train researchers in appropriate use, or monitor compliance. This creates risks of both under-use (missing legitimate benefits) and problematic use (lacking guidance frameworks).

International scholarly organizations and well-resourced institutions bear some responsibility for supporting global capacity development. This might include:

- Developing adaptable frameworks that can be tailored to diverse contexts
- Providing resources for training in local languages
- Supporting research on LLM effects in non-Anglophone and Global South contexts
- Advocating for LLM development that better serves diverse linguistic and cultural communities
- Including Global South perspectives in governance discussions

The governance challenge is genuinely global: LLMs are developed primarily in wealthy countries but affect scholarship worldwide. Inclusive governance requires mechanisms that give voice to affected communities, not merely consultation after decisions are made in Global North contexts.

7. Emerging Policy Responses and Institutional Frameworks

7.1 Publisher Policies

Major publishers have articulated positions on LLM use in submitted manuscripts. While specifics vary, common elements include:

- **Disclosure requirements:** Most major publishers require authors to disclose LLM use in manuscript preparation. Nature Portfolio, Science, Elsevier, and others have implemented such requirements, though specific language varies.
- **Prohibition on LLM authorship:** Publishers consistently maintain that LLMs cannot be listed as authors, reflecting both practical (accountability) and principled (authorship meaning) considerations.
- **Responsibility assignment:** Authors remain fully responsible for all manuscript content, including any LLM-assisted portions. This addresses accountability concerns but places significant verification burdens on authors.
- **Variable specificity:** Some publishers provide detailed guidance on acceptable uses; others offer only general principles. This variation creates challenges for researchers publishing across venues.

7.2 Institutional and Funder Responses

Universities and funding bodies have begun developing policies, though comprehensive frameworks remain uncommon:

- Some institutions have developed guidance documents for researchers
- Others have integrated LLM considerations into research integrity training
- Funding agencies have begun considering disclosure requirements in grant applications and outputs
- A few institutions have created roles or committees specifically addressing AI in research

Institutional responses have generally lagged behind the pace of adoption, creating periods where researchers use LLMs without clear guidance.

7.3 Professional Society Engagement

Disciplinary professional societies play important roles in norm development but have responded unevenly. Some have issued guidance documents; others have convened working groups; many have not yet engaged systematically.

Professional societies may be well-positioned to develop discipline-specific guidance that accounts for the epistemic cultures and practices of particular fields. However, the pace of technological change challenges traditional deliberative processes for developing professional norms.

7.4 Detection, Enforcement, and Their Limits

Efforts to detect LLM-generated text face fundamental challenges. Detection tools produce false positives and false negatives at rates that undermine reliability. As LLMs improve, detection becomes more difficult. The adversarial dynamic—tools designed to evade detection—further complicates matters.

Several implications follow:

First, *enforcement through detection is inherently limited*. Policies that rely primarily on catching LLM use after the fact are unlikely to be effective. This is not a temporary technical problem but a fundamental feature of the detection challenge.

Second, *education and norm development are more promising*. If detection cannot reliably identify problematic use, the alternative is cultivating norms that lead researchers to self-regulate. This requires education about why norms exist, not merely rules to follow.

Third, *verification at multiple stages* may be more effective than post-hoc detection. Structures that build verification into research processes—collaboration, peer feedback during drafting, institutional support for verification—may catch errors regardless of source.

Fourth, *some enforcement mechanisms remain available*. Detection limitations do not mean that egregious cases cannot be identified or that there are no consequences for misconduct. Extreme cases may still be detectable; suspicions may prompt investigation; patterns may emerge. The point is that detection should not be the primary mechanism relied upon.

7.5 Toward Coordinated Governance

Effective governance requires coordination across levels—individual researchers, institutions, publishers, funders, and international bodies. Each level has distinct roles:

Individual researchers bear primary responsibility for appropriate use and verification in their own work.

Institutions can provide training, develop policies, support verification infrastructure, and create cultures that value integrity over productivity metrics that might incentivize problematic LLM use.

Publishers can establish disclosure requirements, develop reviewer guidance, and contribute to norm development through editorial policies and published guidance.

Funders can include LLM-related requirements in grant conditions, support research on LLM effects, and incentivize responsible practices through funding criteria.

International bodies (such as scholarly associations, intergovernmental organizations, and collaborative initiatives) can facilitate coordination across contexts, support capacity development, and advocate for inclusive governance that represents diverse scholarly communities.

Coordination mechanisms might include:

- Multi-stakeholder forums that bring together researchers, institutions, publishers, and funders
- Cross-border initiatives that address the global nature of scholarship
- Iterative policy development that acknowledges uncertainty and allows for revision as evidence accumulates
- Mechanisms for Global South participation in governance decisions

The International Committee of Medical Journal Editors (ICMJE) provides a model: a voluntary association that develops influential guidance adopted broadly across biomedical publishing. Similar cross-publisher coordination for LLM guidance could promote consistency without imposing uniformity.

8. A Framework for Responsible Integration

8.1 Foundational Principles

Building on the foregoing analysis, this section proposes principles for responsible LLM integration in scholarly production:

Principle 1: Transparency

Researchers should disclose LLM use in ways that enable readers to assess potential impacts on the work. Transparency supports informed evaluation and contributes to norm development by making practices visible.

Operationalization: Include statements in methods sections or acknowledgments specifying which LLMs were used, for what tasks, and how outputs were verified. Standardized reporting formats may emerge; in the interim, err toward more disclosure rather than less.

Principle 2: Accountability

Human researchers bear full responsibility for all claims, arguments, and content in their work, regardless of LLM assistance. This principle preserves the accountability structures that underpin scholarly trust.

Operationalization: Before submission, authors should be able to affirm that they understand and endorse all claims in the manuscript and can defend them if challenged. A useful test: Could you explain the reasoning

behind this claim without referring back to the LLM that generated it? If not, the claim may not genuinely be yours.

Principle 3: Verification

LLM outputs should be treated as drafts requiring verification rather than finished products. All factual claims, citations, and substantive assertions should be confirmed through primary sources.

Operationalization: Develop systematic verification workflows. Check all citations against original sources. Verify factual claims through authoritative references. Review analytical outputs for plausibility and correctness. Budget time for verification as part of the research process.

Principle 4: Appropriate Use

LLM applications should be calibrated to context, considering the epistemic requirements of specific tasks, disciplinary norms, and the stakes of potential errors.

Operationalization: The following heuristics may guide judgment:

- Tasks involving factual claims require more stringent verification than stylistic editing
- Novel arguments central to a contribution warrant more human investment than routine descriptions
- High-stakes contexts (clinical recommendations, policy advice) warrant particular caution
- Disciplinary norms about authorship and originality should inform decisions
- Early-career researchers should consider whether LLM use might impede skill development important for their professional growth

Principle 5: Equity Awareness

Researchers, institutions, and policymakers should attend to differential access and impacts, working to ensure that LLM integration does not exacerbate existing inequalities.

Operationalization: Institutions should consider providing equitable LLM access. Researchers should be aware of biases in LLM outputs and actively seek out marginalized perspectives. Policy development should include global perspectives. Individual researchers can advocate for equity considerations in their institutions and professional communities.

Principle 6: Competency Preservation

Researchers and research training programs should maintain practices that develop fundamental scholarly competencies, even where LLMs might substitute for those competencies.

Operationalization: Training programs should ensure that students develop core skills before extensively using LLM assistance. Individual researchers should periodically engage in unassisted practice to maintain skills. Assessment practices should be designed to evaluate core competencies. Mentors should discuss LLM use explicitly with trainees.

8.2 Decision Framework for Specific Use Cases

To operationalize these principles, researchers might employ a structured decision process:

Step 1: Task Characterization

- What type of task is this? (routine/substantive; factual/interpretive; central/peripheral to contribution)
- What are the stakes of potential error?
- What disciplinary norms apply?
- How would colleagues in my field view this use?

Step 2: Capability Assessment

- Am I using LLM assistance because I lack capability, or to increase efficiency?
- If I lack capability, is this a skill I should develop rather than delegate?
- Will I be able to verify the output adequately?
- Am I in a training stage where skill development should take priority?

Step 3: Use Decision

- Is LLM use appropriate for this task in this context?
- If yes, what verification procedures will I employ?
- How will I disclose this use?
- Could I defend this use to colleagues, reviewers, and readers?

Step 4: Verification

- Have I verified all factual claims and citations?
- Do I understand and endorse all content?
- Could I defend this work without reference to the LLM?
- Have I maintained a record of what was LLM-assisted and how it was verified?

8.3 Extended Illustrative Cases

Case 1: Literature Summary for Grant Proposal

Scenario: A researcher uses an LLM to generate summaries of 30 papers for a background section of a grant proposal. The LLM produces fluent summaries.

Analysis: On review, the researcher notices several subtle inaccuracies: papers characterized as finding results they did not find, methodological approaches misdescribed, and one fabricated reference. Correcting these requires returning to primary sources—essentially redoing the literature review.

Outcome: The efficiency gains were minimal after accounting for verification. More importantly, the researcher missed opportunities for insight that close reading might have provided.

Lessons: (a) LLM assistance may not save time when verification requirements are considered. (b) Tasks requiring accuracy may benefit less from LLM assistance than tasks requiring fluency. (c) If you must verify everything anyway, the value of initial LLM generation is reduced.

Case 2: Code Generation for Statistical Analysis

Scenario: A qualitative researcher with limited statistical training receives reviewer requests for quantitative supplementary analyses. She uses an LLM to generate multilevel modeling code, which runs without errors.

Analysis: The researcher cannot fully evaluate whether the model specification is appropriate, whether assumptions are met, or how to interpret variance components. She discloses the LLM assistance but submits the analysis.

Concern: Disclosure does not resolve the competency gap. The analysis may be flawed in ways neither the researcher nor reviewers can identify.

Better approach: The researcher could: (a) collaborate with a quantitative methodologist who can verify the analysis; (b) explicitly acknowledge analytical limitations in the paper; (c) decline to add analysis she cannot fully evaluate, explaining this to editors.

Lessons: (a) The ability to run code does not equal methodological competence. (b) Disclosure is necessary but not sufficient. (c) Collaboration may be more appropriate than LLM-enabled independence for unfamiliar methods.

Case 3: Manuscript Editing for a Non-Native English Speaker

Scenario: A researcher whose first language is Korean uses an LLM to improve grammar and style of a manuscript she has fully drafted. She reviews all edits, reverting those that change meaning.

Analysis: The researcher retains full control over content. The LLM assistance resembles using a copyeditor. The arguments, evidence, and interpretations remain entirely her own.

Outcome: This appears to be legitimate use that preserves authorial accountability while improving accessibility.

Lessons: (a) LLM editing assistance for language polishing raises fewer concerns than substantive content generation. (b) When the researcher retains full understanding and endorsement of all claims, core authorship concerns are addressed. (c) Such use may be particularly valuable for addressing linguistic inequities in global scholarship.

Case 4: Generating a Discussion Section

Scenario: A researcher has completed data analysis but struggles with writing. He uses an LLM to generate a discussion section based on prompts about findings. He edits lightly and submits.

Analysis: The interpretive work of situating findings—connecting them to prior literature, articulating their significance, acknowledging limitations, suggesting implications—is a core authorial contribution. If the researcher cannot articulate why certain connections are made or what follows from the findings, genuine authorship may be compromised.

Concern: The discussion represents the researcher's interpretation of findings. If it substantially represents LLM interpretation rather than researcher interpretation, something essential to authorship is missing.

Better approach: The researcher might: (a) use the LLM output as a starting point for his own thinking, substantially revising to express his genuine interpretation; (b) develop outlining skills with LLM assistance but write the actual prose; (c) dictate his thoughts and have the LLM refine prose rather than generate arguments.

Lessons: (a) Interpretive work is not readily delegable without compromising authorship. (b) The distinction between stylistic assistance and substantive generation matters. (c) If you could not articulate the discussion's claims without reference to the LLM, something has gone wrong.

Case 5: Historical Research with Archival Sources

Scenario: A historian working on 18th-century French colonial administration uses an LLM to translate archival documents from French and to summarize secondary literature in German that she does not read fluently.

Analysis: Translation assistance enables engagement with sources otherwise inaccessible—broadening the evidentiary base. However, LLM translations may contain errors that affect interpretation. The historian cannot fully verify German translations she cannot read.

Approach: For French translations, she verifies by reading originals (using LLM to speed initial processing). For German sources, she: (a) treats LLM summaries as preliminary, seeking collaboration with German-reading colleagues for key sources; (b) explicitly acknowledges the mediated nature of her engagement with German scholarship; (c) focuses interpretive weight on sources she can verify directly.

Lessons: (a) LLM translation assistance can broaden access but introduces verification challenges. (b) Transparency about mediated access is important. (c) Epistemic weight should track verification capacity.

8.4 Addressing Conflicts Between Principles

The principles proposed may sometimes tension with each other:

- **Transparency vs. Stigma:** Full disclosure might invite unfair stigma if LLM use is viewed negatively, even when use is appropriate. *Resolution:* As norms develop, transparency helps establish that appropriate use is legitimate. Individual risk of stigma must be weighed against collective benefit of norm development. Researchers might advocate for fair evaluation standards alongside practicing disclosure.
- **Equity (democratization) vs. Competency Preservation:** Broad LLM access might enable less-trained individuals to produce sophisticated-seeming work, potentially undermining incentives for skill development. *Resolution:* Access should be accompanied by training in appropriate use. The goal is augmentation after foundational development, not substitution for it. Institutions bear responsibility for training, not just access.
- **Efficiency vs. Verification Burden:** LLMs promise efficiency, but verification requirements may offset gains. *Resolution:* Accept that some uses yield less efficiency gain than expected once verification is accounted for. Calibrate LLM use to tasks where verification is manageable relative to gains.
- **Appropriate Use vs. Individual Discretion:** What counts as appropriate varies by context, making enforcement difficult. *Resolution:* Rely primarily on professional norms, education, and communities of practice rather than external policing. Accept that judgment calls will differ; focus on egregious cases rather than attempting to regulate all decisions.
- **Innovation vs. Precaution:** Restrictive approaches might prevent beneficial innovation; permissive approaches might enable harm. *Resolution:* Adopt iterative, evidence-responsive frameworks. Begin with caution in high-stakes contexts; gather evidence; adjust as understanding deepens. Avoid both paralysis and recklessness.

9. Temporal Dynamics and Future Trajectories

9.1 Evolving Practices and Normalization

Current patterns likely represent early-stage adoption. As familiarity increases, practices may evolve in several directions:

- **Normalization:** LLM assistance may become routinely integrated and less remarked upon, much as word processors or citation managers are not typically disclosed. Whether this normalization is desirable depends on whether the concerns identified above prove serious or manageable.
- **Specialization:** Purpose-built LLMs for specific disciplines or tasks may emerge, offering more reliable domain-specific assistance. Such specialization might address some concerns about generic LLM limitations while raising new questions about disciplinary capture.
- **Skill adaptation:** Researchers may develop new competencies around effective LLM prompting, verification, and integration. These meta-competencies may become part of what it means to be a skilled researcher—not replacing traditional competencies but supplementing them.
- **Generational effects:** Researchers trained with LLMs from the start of their careers may develop different relationships to these tools than those who adopted them later. The long-term effects of different developmental trajectories are genuinely unknown.

9.2 Recursive Effects and Training Data Dynamics

As LLM-assisted scholarship enters the published literature and becomes training data for future models, recursive effects may emerge that warrant monitoring:

- **Amplification of patterns:** Tendencies in LLM-assisted writing may become more pronounced as models train on LLM-influenced text. Stylistic homogenization might compound.
- **Error propagation:** Inaccuracies introduced through LLM use might propagate if incorporated into training data without correction. The literature might develop persistent errors that become harder to correct.
- **Epistemic feedback loops:** If LLMs influence what research questions are pursued, what framings are employed, and what conclusions are drawn, and if LLM-influenced scholarship then trains future LLMs, a feedback loop emerges that could shape the direction of knowledge in ways that are difficult to trace or interrupt.
- **Citation pattern effects:** If LLMs reinforce existing citation patterns (citing heavily-cited work), the Matthew Effect in scholarship might intensify, further marginalizing less-visible work.

These dynamics suggest the importance of: (a) curating training data to maintain quality and diversity; (b) preserving LLM-independent scholarship as a check and comparison; (c) monitoring for signs of problematic convergence; (d) maintaining human judgment as a counterweight to LLM influence.

9.3 Resistance, Alternatives, and Plural Futures

Not all scholars or communities will adopt LLMs. Some may resist on principled grounds—concerns about authenticity, craft traditions, epistemic values, or broader critiques of AI technology. Such resistance should not be dismissed as mere technophobia; it may reflect considered judgments about what particular scholarly communities value.

Possible manifestations of resistance include:

- Disciplinary or subdisciplinary norms limiting LLM use
- Journals or conferences that explicitly prohibit or discourage LLM assistance
- Research groups or traditions that emphasize LLM-free practice as part of their identity
- Individual scholars who choose not to use LLMs for philosophical or practical reasons

The coexistence of LLM-intensive and LLM-resistant communities creates interesting dynamics. Will work produced without LLM assistance be valued differently—perhaps seen as more authentic, or perhaps as inefficiently produced? Will there be "LLM-free" certifications or venues? How will mixed collaborative teams navigate different practices?

These questions do not have predetermined answers. Technological adoption is not inevitable or uniform; communities make choices that shape technological trajectories. The future will likely be plural—different communities adopting different practices for different reasons.

9.4 Scenarios for Future Development

Several scenarios for longer-term development warrant consideration:

1. **Scenario 1: Seamless integration.** LLMs become normalized as writing and research tools, integrated into scholarly workflows much as word processors were. Concerns about skill atrophy prove exaggerated; new competencies develop; scholarship continues largely as before, with enhanced efficiency.
2. **Scenario 2: Stratified adoption.** Wealthy institutions integrate LLMs extensively; resource-poor contexts lag behind. Inequality in scholarly production increases. Global knowledge hierarchies intensify.
3. **Scenario 3: Quality erosion.** LLM integration enables productivity but undermines quality. Verification fails to keep pace with generation. The scholarly literature becomes less reliable; trust in scholarship declines.
4. **Scenario 4: Reconceptualization.** LLM integration prompts fundamental rethinking of authorship, originality, and expertise. New concepts and practices emerge that are neither traditional scholarship nor mere LLM output but something genuinely new.
5. **Scenario 5: Backlash and retrenchment.** High-profile failures or scandals involving LLM use prompt restrictive policies. A chilling effect reduces adoption below optimal levels.

None of these scenarios is predetermined. Outcomes will depend on choices made by researchers, institutions, publishers, and policymakers—choices that the present analysis aims to inform.

10. Conclusion

The integration of Large Language Models into scholarly production represents a development whose implications extend beyond questions of efficiency and appropriate use to fundamental transformations in how knowledge is produced, validated, and understood. This article has examined LLM capabilities and applications, analyzed implications for epistemic quality and research integrity, addressed equity concerns with particular attention to global dimensions, surveyed emerging policy responses, and proposed a framework for responsible integration.

Several key insights emerge from this analysis:

First, LLMs are not merely neutral tools but participants in networks of knowledge production that may reshape practices in ways requiring careful examination. The STS frameworks employed reveal that LLM integration involves not just governance of new technology but potential reconfiguration of foundational categories—authorship, originality, expertise.

Second, the implications of LLM integration vary substantially across disciplines, reflecting different epistemic cultures and practices. The detailed case studies demonstrate that what counts as appropriate use depends on what a discipline values and how it produces knowledge.

Third, benefits including enhanced accessibility and efficiency must be weighed against risks including factual unreliability, bias propagation, and potential skill atrophy. The evidence base for assessing these tradeoffs remains limited; many claims represent theoretical arguments rather than established findings.

Fourth, equity considerations—particularly regarding global access, linguistic justice, and the representation of diverse knowledge traditions—deserve central attention. LLM integration risks intensifying existing inequalities unless explicitly designed to counter them.

Fifth, effective governance requires coordination across individuals, institutions, publishers, funders, and international bodies. Detection-based enforcement is inherently limited; education, norm development, and structural support for verification are more promising.

Sixth, the future is not predetermined. Different scholarly communities may adopt different practices; resistance to LLM integration may prove as significant as adoption. The choices made during this formative period will shape trajectories for decades.

The framework proposed—emphasizing transparency, accountability, verification, appropriate use, equity awareness, and competency preservation—provides guidance while acknowledging that guidance itself may require revision as understanding deepens. The operationalized decision procedures and extended case studies offer practical resources. The attention to coordinated governance suggests mechanisms for collective action.

Important limitations should be acknowledged. The empirical evidence base remains limited; many claims are necessarily speculative. Philosophical questions about machine cognition remain contested. The analysis, despite efforts to incorporate diverse perspectives, reflects the author's position within particular scholarly traditions. The practical recommendations may require adaptation as technology evolves.

The academic community's response to LLMs will shape the future of scholarly production. Thoughtful integration that leverages genuine benefits while safeguarding epistemic values can strengthen research practices. However, this requires neither uncritical adoption nor reflexive rejection but sustained, critical engagement that attends to disciplinary specificity, global diversity, temporal dynamics, and the fundamental questions about scholarship that LLMs bring to the fore.

Perhaps most importantly, this moment of technological change provides an occasion for reflection on what scholarship is for—what values it serves, what practices sustain it, and what futures we wish to create. These questions are not merely about managing new technology but about the nature and purpose of knowledge production. Engaging them well requires the very capacities—careful reasoning, attentive reading, original thinking, and collaborative inquiry—that scholarship at its best has always cultivated.

References

- Barthes, R. (1967). The death of the author. *Aspen*, 5–6.
- Beigel, F. (2014). Introduction: Current tensions and trends in the World Scientific System. *Current Sociology*, 62(5), 617–625.
- Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., ... Amodei, D. (2020). Language models are few-shot learners. In *Advances in Neural Information Processing Systems* (Vol. 33, pp. 1877–1901). Curran Associates.
- Callon, M. (1986). Some elements of a sociology of translation: Domestication of the scallops and the fishermen of St Brieuc Bay. In J. Law (Ed.), *Power, action and belief: A new sociology of knowledge?* (pp. 196–223). Routledge.
- Canagarajah, A. S. (2002). *A geopolitics of academic writing*. University of Pittsburgh Press.
- Checco, A., Bracciale, L., Loreti, P., Pinfield, S., & Bianchi, G. (2021). AI-assisted peer review. *Humanities and Social Sciences Communications*, 8(1), Article 25.
- Chen, M., Tworek, J., Jun, H., Yuan, Q., Pinto, H. P. de O., Kaplan, J., Edwards, H., Burda, Y., Joseph, N., Brockman, G., Ray, A., Puri, R., Krueger, G., Petrov, M., Khlaaf, H., Sastry, G., Mishkin, P., Chan, B., Gray, S., ... Zaremba, W. (2021). Evaluating large language models trained on code. *arXiv preprint*. arXiv:2107.03374
- Clark, A. (2008). *Supersizing the mind: Embodiment, action, and cognitive extension*. Oxford University Press.
- Cole, D. (2020). The Chinese Room Argument. In E. N. Zalta (Ed.), *The Stanford Encyclopedia of Philosophy* (Winter 2020 Edition). Stanford University.
- Dennett, D. C. (1991). *Consciousness explained*. Little, Brown and Company.
- Ferrara, E. (2023). Should ChatGPT be biased? Challenges and risks of bias in large language models. *First Monday*, 28(11).
- Fischman, G. E., & Alperin, J. P. (2015). Visibility and quality in Spanish-language Latin American scholarly publishing. *Information Technologies & International Development*, 11(4), 1–21.
- Flowerdew, J. (2019). The linguistic disadvantage of scholars who write in English as an additional language: Myth or reality. *Language Teaching*, 52(2), 249–260.
- Foucault, M. (1984). What is an author? In P. Rabinow (Ed.), *The Foucault reader* (pp. 101–120). Pantheon Books.
- Fricker, M. (2007). *Epistemic injustice: Power and the ethics of knowing*. Oxford University Press.
- Gao, C. A., Howard, F. M., Markov, N. S., Dyer, E. C., Ramesh, S., Luo, Y., & Pearson, A. T. (2022). Comparing scientific abstracts generated by ChatGPT to original abstracts using an artificial intelligence output detector, plagiarism detector, and blinded human reviewers. *bioRxiv preprint*. <https://doi.org/10.1101/2022.12.23.521610>
- Hyland, K. (2004). *Disciplinary discourses: Social interactions in academic writing*. University of Michigan Press.
- Ihde, D. (1990). *Technology and the lifeworld: From garden to earth*. Indiana University Press.
- Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y. J., Madotto, A., & Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), Article 248.
- Knorr-Cetina, K. (1999). *Epistemic cultures: How the sciences make knowledge*. Harvard University Press.
- Latour, B. (2005). *Reassembling the social: An introduction to actor-network-theory*. Oxford University Press.

- Lund, B. D., & Wang, T. (2023). Chatting about ChatGPT: How may AI and GPT impact academia and libraries? *Library Hi Tech News*, 40(3), 26–29.
- Pickering, A. (1995). *The mangle of practice: Time, agency, and science*. University of Chicago Press.
- Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417–424.
- Simkin, M. V., & Roychowdhury, V. P. (2003). Read before you cite! *Complex Systems*, 14(3), 269–274.
- Star, S. L., & Griesemer, J. R. (1989). Institutional ecology, 'translations' and boundary objects: Amateurs and professionals in Berkeley's Museum of Vertebrate Zoology, 1907-39. *Social Studies of Science*, 19(3), 387–420.
- Stokel-Walker, C., & Van Noorden, R. (2023). What ChatGPT and generative AI mean for science. *Nature*, 614(7947), 214–216.
- Thelwall, M. (2022). Can the quality of published academic journal articles be assessed with machine learning? *Quantitative Science Studies*, 3(1), 208–226.
- Van Noorden, R., & Perkel, J. M. (2023). AI and science: What 1,600 researchers think. *Nature*, 621(7980), 672–675.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. In *Advances in Neural Information Processing Systems* (Vol. 30, pp. 5998–6008). Curran Associates.
- Wei, J., Tay, Y., Bommasani, R., Raffel, C., Zoph, B., Borgeaud, S., Yogatama, D., Bosma, M., Zhou, D., Metzler, D., Chi, E. H., Hashimoto, T., Vinyals, O., Liang, P., Dean, J., & Fedus, W. (2022). Emergent abilities of large language models. *Transactions on Machine Learning Research*.