

Level 2: Applied

Practicing Librarian

## Module 06

# AI for research support

Research support is where many librarians first encounter real AI utility, and real AI failure. This module maps where AI adds value to literature reviews, search strategy, and evidence synthesis, and where it gets students into trouble.

**By Yulia Brusova**

25 min read · Reviewed June 2026

---

**AI for Academic Libraries** · [ai-in-academic-libraries.vercel.app](https://ai-in-academic-libraries.vercel.app)

© 2026 · Licensed under CC BY-NC-SA 4.0

#### WHAT YOU'LL BE ABLE TO DO

- 1 Use AI to draft, refine, and expand a database search strategy
- 2 Apply AI for summarizing and synthesizing research documents patrons provide
- 3 Describe at least three AI native research tools and when to recommend them
- 4 Explain to a patron the appropriate and inappropriate uses of AI in their research process
- 5 Develop a research support workflow that integrates AI at appropriate points

*A pattern that appears repeatedly across disciplines and student levels is students arriving at the reference desk with a list of twenty AI generated sources for their literature review. In a typical such list, eight articles do not exist, four exist but say nothing resembling what the student claims they say, and the remaining eight require individual verification before they can be used. This experience defines the challenge at the center of this module. It is not that AI is a poor tool for research support. It is that the workflows most students naturally apply to AI in a research context are poor workflows: they produce the appearance of scholarship without the substance of it. In order to provide effective research support in an AI environment, librarians need to know precisely where AI adds genuine value, where it introduces serious risk, and how to communicate that distinction in ways that change how patrons actually work.*

## Where AI genuinely helps in research support

The distinction that matters most for research support is between AI tasks grounded in content the patron has already retrieved versus AI tasks that require generating facts from training data. In research support contexts, the applications where AI delivers genuine value tend to be the former: assisting patrons with thinking through a research problem, processing documents they already hold, and explaining concepts they have not yet encountered. Understanding this distinction allows librarians to direct patrons toward productive AI use without endorsing the problematic applications that lead to fabricated citations.

**Search strategy development** is the area of most consistent genuine value. When a patron brings a research question and needs help developing a search strategy, AI can rapidly generate a substantive set of search terms, synonyms, related concepts, Boolean combinations, and relevant subject headings. For example, a librarian working with a

graduate student researching the intersections of telehealth adoption and rural health equity can ask Claude to generate search terms organized by concept cluster (telehealth terminology, rural health identifiers, equity and access language, and relevant MeSH headings for PubMed) in approximately thirty seconds. Such a brainstorm does not replace the librarian's judgment about which terms are most appropriate for which databases; it provides the raw material for that judgment with substantially less effort than generating it manually.

**Summarizing provided documents** is the second reliably useful application. When a patron has a set of abstracts, a full article, or a selection of PDFs and needs help identifying key themes, methodological approaches, or relevant findings, AI working with that provided content produces reliable and useful summaries. This is the information in mode discussed in Module 02: the model is processing text the patron has already retrieved from authoritative sources rather than generating facts from training data. For example, a nursing student preparing a systematic review can paste twelve abstracts into Claude and ask it to identify shared themes, divergent findings, and methodological differences across the set, and receive a structured synthesis that would take considerably longer to produce manually. Such summaries warrant review rather than acceptance without evaluation, but the hallucination risk is substantially lower than when AI is asked to generate content from scratch.

**Explaining research concepts** on demand is a third consistent strength. Research methods, statistical approaches, citation conventions, and disciplinary norms that a patron encounters for the first time can be explained at the appropriate level by AI without requiring an available librarian. For example, a social work student unfamiliar with the difference between a systematic review and a scoping review can receive a clear, example grounded explanation in plain language immediately. Such on demand explanation is particularly valuable in research support contexts outside of library hours or when a patron is working independently between consultations, with the important caveat that any substantive claims the explanation contains should be confirmed against authoritative methodology documentation before being relied upon.

Additionally, AI serves research support well at the **question formation** stage. A patron struggling to articulate a focused research question can ask AI to generate five alternative framings of a broad topic, or to suggest what aspects of a topic appear underresearched. While the model's knowledge of any given field has limitations, including a training cutoff date and representation biases toward English language and Western sources, it can reliably generate the conceptual variety that helps patrons move from a vague interest to a workable question. Such assistance is often invisible as "research support," but it represents a

meaningful expansion of the librarian's consultative role when AI is positioned as a thinking partner in the consultation rather than a source of authoritative answers.

## **Where AI fails: the research support risks librarians must communicate**

The failure modes of AI in research support are structural, not incidental. They follow directly from how large language models work, specifically prediction from training data, and they do not disappear with newer model versions or more careful prompting. In order to advise patrons effectively, librarians need to understand these failure modes at a level that allows them to explain the underlying reason, not merely list a set of prohibitions. Patrons who understand why AI fails in specific ways develop more durable verification habits than patrons who have simply been warned.

**Finding sources** is the most consequential failure mode. AI cannot search databases. When asked to find peer reviewed articles on a topic, the model generates seemingly credible citations from statistical patterns in its training data. Some of these citations correspond to real publications; many do not. Such fabrications are not random; they follow the conventions of real citations closely, including realistic author names, appropriate journal titles, and plausible publication years, which makes them difficult to identify without verification. For example, a history student who asks ChatGPT for primary sources on the 1930s labor movement may receive a list that includes real archives, fabricated archive names, and citation details that are internally plausible but externally unverifiable. The danger is not that the citations look wrong; it is that they look right.

**Current literature** presents a related problem. Most major AI tools have a training cutoff date, a point after which no new information was incorporated, which ranges from one to several years behind the present. For patrons researching fast moving topics such as clinical trials, pharmaceutical approvals, policy changes, or recent legislative developments, AI is actively misleading: the model responds with training data confidence about a landscape that may have changed substantially. For example, a policy student researching current federal AI regulation who asks Claude for an overview will receive information current at the model's training cutoff, presented without any indication of how rapidly the regulatory environment has shifted since. Such outdated information is particularly dangerous because it arrives in the same confident, well organized prose as accurate information.

**Specialized and niche subjects** are underrepresented in AI training data in proportion to how sparsely they were written about during the training period. This affects research support

unevenly: a librarian helping a patron with a mainstream topic in psychology or history can reasonably expect more complete training data coverage than a librarian supporting research on an emergent subfield, a non English language scholarship area, or a topic whose primary literature exists in specialized professional databases rather than the general web. For example, a patron researching Indigenous language preservation efforts in a specific region may find that AI synthesizes primarily from a small number of English language sources, producing a response that reflects the anglophone framing of the subject rather than the community centered scholarship that would be most appropriate. Librarians should treat AI coverage of any specialized or underrepresented topic as presumptively incomplete.

**Evaluating source quality** is a task AI cannot perform reliably. A patron who asks AI to identify the most credible or methodologically rigorous sources from a list is asking the model to apply a judgment it has no mechanism to exercise. AI evaluates sources the same way it evaluates everything, by pattern matching to what sounds authoritative, which means it may identify a source as credible based on prestige vocabulary in an abstract rather than on any actual assessment of methodology, peer review status, or scholarly standing. In order to protect patrons from this misunderstanding, the instruction point must be explicit: AI can describe what a source claims; it cannot assess whether those claims are warranted. Source evaluation remains a library professional competency, and it remains necessary regardless of how sophisticated AI responses become.

## **Search strategy development: using AI to think through the search, not conduct it**

Search strategy development is the research support application where AI most clearly earns its place in the consultation workflow, and understanding precisely how it contributes, and precisely where it stops contributing, is the practical knowledge this section develops. The key distinction, which bears repeating to patrons as well as colleagues, is that AI assists with thinking about how to search, not with searching itself. It generates terms, structures, and conceptual frameworks; it does not retrieve records, query databases, or access controlled vocabularies in real time.

The core workflow proceeds in stages that the librarian can model explicitly in a research consultation. First, articulate the research question clearly: the more specific the question the patron provides, the more useful the AI generated search structure will be. For example, a question stated as "the effects of nurse to patient ratios on patient outcomes in intensive care settings" will produce a more targeted set of search terms than "nursing and hospitals." Second, ask AI to generate search terms organized by concept cluster: "Generate search

terms for each concept in this research question, including synonyms, related terms, alternate spellings, and relevant controlled vocabulary for PubMed." Such a prompt typically produces a structured set of term clusters that the librarian and patron can review and refine together.

Third, ask AI to suggest Boolean combinations and search string structures. For example, prompting Claude with "Suggest three Boolean search strings using these term clusters, with AND connecting the concepts and OR connecting synonyms within each cluster" produces ready-to-test strings. The librarian's professional role is to evaluate these strings against the specific database's indexing conventions: MeSH terms for PubMed, Thesaurus descriptors for CINAHL, controlled vocabulary for PsycINFO, and adjust accordingly. AI can suggest that a medical search should include MeSH terms; it cannot guarantee that the specific terms it names are current and correctly formatted. Such verification is the librarian's contribution, not an optional addition.

Fourth, iterate in response to database results. When an initial search produces too many or too few results, or retrieves results in the wrong subarea of the topic, the patron can describe this to AI and ask for revised strategies: "My search for these terms returned 3,000 results in CINAHL, most of which are not focused on the intensive care setting. Suggest three strategies for narrowing: adding limiters, adding additional concept terms, or restructuring the Boolean." Such iteration makes the consultation more dynamic and often surfaces search angles that neither the librarian nor the patron had considered independently. It is useful to think of AI as a well read colleague who has processed an enormous quantity of methodology literature and library instruction content; that colleague can generate options quickly, but the librarian's knowledge of specific database indexing is what makes those options actionable.

One additional application worth making explicit: AI can help patrons develop search strategies for databases they have never used, by explaining how a database is structured, what types of sources it covers, and how its search interface differs from familiar tools. For example, a social science student accustomed to Google Scholar who is being introduced to Sociological Abstracts for the first time can ask AI to explain the database's scope, subject heading structure, and most effective search approaches for the discipline. Such orientation reduces the learning curve for new database users in ways that complement rather than replace library instruction, and it positions AI as a preparation tool for library work rather than a replacement for it.

## Summarizing and synthesizing research documents patrons provide

The information in mode of AI use (providing documents for AI to process rather than asking AI to generate information from training data) is the highest reliability application in research support contexts, and it is worth developing as a core consultation tool rather than treating it as an incidental capability. When a patron brings documents to an AI interaction, such as abstracts, full articles, PDFs, or reports, and asks AI to help process that content, the hallucination risk drops substantially because the model is working with provided text rather than generating facts from training. Such a distinction should inform how librarians introduce AI use to patrons at the research support stage: emphasize working with sources the patron has already verified over asking AI to generate sources from scratch.

The most immediate application is abstract triage. A patron who has run a database search and retrieved thirty abstracts faces a time intensive reading task, working through all thirty to identify which articles are worth pursuing to full text, which are tangentially related, and which are not relevant at all. AI can perform an initial triage pass by processing the full set of abstracts and organizing them by relevance, methodology, or research focus. For example, a graduate student can paste thirty PubMed abstracts into Claude and ask: "Organize these abstracts into three groups: highly relevant to a systematic review on this topic, moderately relevant and worth reviewing in full text, and tangentially related. For the highly relevant group, note what each one contributes." Such triage reduces a forty five minute reading task to a ten minute review, freeing the patron's time for deeper engagement with the most relevant sources.

For more advanced synthesis tasks (identifying shared themes across a literature, mapping methodological approaches, or characterizing the evidence base for a specific intervention), AI working with provided documents can produce structured syntheses that support rather than replace the patron's own intellectual engagement. In order to use this capability responsibly, librarians should explain to patrons that AI synthesis of provided documents is a starting point for their own analysis, not a finished product. For example, a public health student asking AI to synthesize the methodology sections of twelve provided articles may receive a useful overview of study designs, sample sizes, and outcome measures, but the student must engage with that synthesis critically and confirm key details against the original documents before incorporating it into a literature review.

The limit of this approach is important to communicate explicitly. AI processes provided text and reflects that content back in organized form; it does not have access to the broader

literature and cannot tell the patron what the provided documents omit. A patron who synthesizes fifteen articles and concludes that the literature "shows X" has synthesized fifteen articles, not the field. For example, if all fifteen articles happen to represent one methodological tradition or one disciplinary perspective, the synthesis will reflect that perspective without noting its limitations. The librarian's role (identifying that a broader literature search should precede or accompany AI synthesis, and that the patron's selection criteria shape what the synthesis represents) is not displaced by AI's facility with provided documents. Indeed, it becomes more important, because the ease of AI synthesis can create an illusion of comprehensiveness that a narrowly selected document set does not warrant.

## **AI native research tools: a librarian's evaluation**

General purpose AI tools such as ChatGPT, Claude, and Gemini were designed for broad conversational use and carry the hallucination risks in research contexts described in this module's earlier sections. A distinct category of tools has been built specifically to address those risks by grounding AI responses in actual database searches or controlled scholarly corpora. These tools represent a meaningful improvement over general purpose AI for certain research support tasks, and librarians who understand their capabilities and limitations can recommend them credibly to patrons who need AI assistance in contexts where fabricated citations would be professionally or academically costly.

**Connected Papers** is an AI enhanced discovery tool that visualizes the citation network surrounding a specific paper. A patron enters a known relevant article, and the tool maps both the papers that cite it and the papers it cites, producing a visual graph that makes the intellectual neighborhood of a topic navigable. For example, a sociology graduate student who has identified one foundational paper on their topic can use Connected Papers to surface the cluster of scholarship that has engaged with that paper, discover articles they had not found through keyword searching, and map the structure of the scholarly conversation in their field. Such citation network visualization is particularly valuable in the early stages of a literature review, when a patron needs to understand the landscape of a field before committing to a specific research angle. Connected Papers does not generate text from training data; it retrieves from real citation networks, which means it carries no hallucination risk for bibliographic purposes. It shows connections among documents that actually exist.

**Elicit** is an AI research assistant that searches Semantic Scholar, a database of over 200 million academic papers, and returns real papers with AI generated summaries drawn from the actual document content. For example, a clinical librarian helping a patron with a

systematic review can use Elicit to run structured searches, organize results by study design, and generate comparative summaries across a retrieved set of papers. The distinction from general purpose AI is fundamental: Elicit grounds its responses in real search results from a controlled scholarly database, which substantially reduces hallucination risk for bibliographic content. Such grounding does not eliminate all risk; AI summaries of real papers can still mischaracterize nuanced findings, and Semantic Scholar's coverage is not equivalent to the major licensed research databases, but it is a considerably more reliable starting point for literature review support than asking ChatGPT to find sources. Elicit is most useful for systematic review workflows, rapid literature mapping, and extracting specific data points across a large set of papers.

**Consensus** takes a complementary approach: it searches peer reviewed papers and attempts to synthesize the overall state of evidence on empirical questions. For example, a patron asking "Does sleep deprivation affect academic performance?" receives a structured response indicating how many papers Consensus found, how the evidence divides across positive and null findings, and links to the actual papers supporting each position. Such evidence mapping is most useful for quick evidence check tasks where the patron needs a rapid orientation to the state of the research before pursuing a deeper literature review. The free tier is genuinely usable for most research support purposes, and the tool handles empirical questions in medicine, psychology, and social science well. It is less effective for humanities topics, historical questions, and highly specialized or emergent research areas that are not well represented in its corpus.

**Perplexity in academic mode** combines AI synthesis with real time web search and includes citations for every response. For research support, it is useful for factual orientation, helping a patron understand the policy context for a research topic, identifying key organizations in a field, or getting a quick overview of a methodology, because the citations allow the patron to verify the AI's characterizations against primary sources. Such verification enabled use is more appropriate than general AI use for the orientation stages of research, though Perplexity does not replace database searching for peer reviewed literature and its summaries carry the same risk of mischaracterization as other AI generated content.

In order to recommend any of these tools to patrons with confidence, librarians should experiment with them personally on topics they know well, so that they can evaluate the accuracy of summaries and the quality of retrieved results before directing patrons to rely on them. Such personal evaluation, including assessing whether Elicit's summaries of articles the librarian has already read are accurate, produces the kind of calibrated, firsthand judgment that credible patron recommendations require.

A second category of AI native tool operates not over an external corpus like Semantic Scholar but over the databases and discovery layers the library already provides access to, and this category has moved into production over the past year. Ex Libris's Primo Research Assistant performs retrieval augmented generation over the Central Discovery Index and returns a small set of the most relevant sources with citations back to records the patron can open directly; Clarivate offers a parallel Summon Research Assistant for institutions on that discovery platform. JSTOR, EBSCO (as AI Insights), Ebook Central, and Scopus have each added AI assistants that operate over their own indexed content, generating summaries and suggested sources that are traceable to records the library's subscriptions actually cover. For example, a patron working within EBSCO who asks AI Insights to summarize an article's argument receives a summary grounded in that specific article, not in a language model's general training data, which is precisely the distinction this module has emphasized between AI working with provided or retrieved content and AI generating from training data alone. In order to recommend these tools confidently, librarians should check which of them are active under their institution's specific subscriptions, since rollout has varied by vendor and license tier; where available, they belong alongside Elicit, Connected Papers, and Consensus, with the additional advantage that they operate inside resources the library has already licensed and the patron is already authenticated into.

A further development worth knowing about, even where it is not yet in routine use, is Elsevier's Deep Research feature within Scopus: an agentic research tool grounded in the Scopus database that, rather than returning a single response, plans and executes a multistep research process and displays its reasoning on screen as it works, namely which searches it is running, which sources it is consulting, and how it is narrowing toward a synthesis. Module 01 introduced the vocabulary for this category, agentic AI, systems that set goals, plan tasks, and act with minimal guidance, and Deep Research is a concrete instance of that vocabulary applied to research support. For a patron, watching an agentic tool's reasoning unfold on screen is a different experience than receiving a finished answer, because each step, each search run, each source consulted, is visible and could in principle be checked. Such visibility does not mean the output requires less verification; a multistep process grounded in Scopus is still a process whose final synthesis the patron must evaluate using the same judgment this module has built throughout. What changes is the texture of the interaction, not the professional standard applied to it. Whatever degree of autonomy these tools reach, the determination of whether a synthesis actually answers the patron's question, represents the literature fairly, and is ready to inform their work remains a judgment that belongs to the patron and the librarian supporting them, not to the tool that produced it.

## Building a research support workflow that integrates AI at appropriate points

The most effective way to communicate AI's role in the research process to patrons is not a list of rules but a map of the research workflow with AI's appropriate and inappropriate uses indicated at each stage. Such a map makes the distinction concrete and contextual rather than abstract, and it can be developed as a shareable handout, a consultation framework, or an instructional tool depending on the librarian's setting and patron population.

A research workflow for most undergraduate and graduate work proceeds through recognizable stages: **topic development and question formation**, **literature search and retrieval**, **source evaluation**, **reading and synthesis**, and **writing and citation**. AI's role differs at each stage, and the appropriate distinctions are not difficult to communicate clearly once they have been mapped explicitly.

At the **topic development** stage, AI is genuinely useful. Exploring a topic's scope, generating potential research angles, identifying related fields and interdisciplinary connections, and developing a focused research question from a broad interest are all tasks where AI's brainstorming capabilities serve the patron without requiring verified factual claims. For example, a first year student with a broad interest in climate change can use AI to narrow toward a specific angle (the intersection of climate change and food security in sub Saharan agriculture) before bringing that focused question to the library for database searching. Such use is entirely appropriate and represents a net gain in research preparation.

At the **literature search and retrieval** stage, AI's appropriate role is limited to search strategy development, specifically generating term clusters, Boolean combinations, and database recommendations, and does not extend to finding sources. The student brings the search strategy to actual databases: JSTOR, PubMed, PsycINFO, or whatever the appropriate disciplinary resource is. For example, the workflow instruction for this stage can be stated simply: "Use AI to develop your search terms; use the library database to find the sources." Such a clear two step instruction prevents the most common and consequential research workflow error.

At the **source evaluation** stage, AI should not be used for quality assessment, for the reasons described earlier in this module. The patron must evaluate sources using conventional information literacy criteria (authority, accuracy, currency, coverage, and purpose) applied to the source itself. Such evaluation requires direct engagement with the source and cannot be outsourced to a model that cannot access the actual document or assess its scholarly standing.

At the **reading and synthesis** stage, AI is most useful when the patron has already retrieved and verified the sources. Summarizing an abstract, explaining a methodology section in plain language, identifying the key argument of a dense theoretical paper, or synthesizing themes across a provided set of abstracts are all appropriate uses that work with documents the patron already holds. For example, a student who has retrieved twelve legitimate articles and needs to organize them for a literature review can use AI to produce an initial thematic map from the abstracts, which the student then reviews, adjusts, and develops with their own analytical engagement.

At the **writing and citation** stage, the primary guidance is that AI generated citations must never appear in final work without independent verification, and that AI assistance with prose should be disclosed according to the institution's and instructor's policies. In order to prevent citation fabrication at this stage, the workflow should treat verification as a required step rather than an optional precaution: before any citation is included in a paper, the patron must locate the actual document in a library database. Developing this map as a printable or digital handout for research consultations is among the most immediately practical steps a librarian can take to improve research AI literacy, because a workflow framework gives patrons something concrete to follow, rather than a prohibition to work around.

## **Your expanded instruction role: research AI literacy as a core library competency**

Reference and instruction librarians have always mediated between patrons and information sources that carry specific risks: teaching database searching to users who default to Google, explaining the distinction between popular and scholarly sources to students who cannot recognize it, and building citation verification habits in communities where cutting corners is normalized. AI in research contexts presents the same professional challenge: a powerful tool, widely adopted by patrons before its failure modes are understood, creating predictable errors that the library is positioned to prevent.

What is new is not the challenge but the scale and the speed. Students and faculty are adopting AI research workflows faster than institutions can develop formal guidance, which means librarians are encountering AI related research errors in reference consultations, embedded instruction sessions, and research help desk tickets before any institutional response is in place. In order to address this effectively, librarians need a proactive instruction approach rather than a reactive correction approach: teaching appropriate research AI use as part of standard library instruction, not as a specialized add on for students who are already in trouble.

The instruction content that matters most at the research support level falls into three categories. First, **workflow instruction**: explicitly teaching patrons where AI fits and does not fit in the research process, using the workflow map described in the prior section. This is most effective when delivered in the context of a specific assignment rather than as a general AI literacy session, because contextual specificity produces more durable behavior change than abstract instruction. For example, a single shot session for a research methods course can include five minutes on the appropriate AI workflow for the literature review required in that course, addressing this specific assignment, this database, and this citation requirement rather than AI in general.

Second, **live demonstration**: showing students what AI citation fabrication looks like by attempting to locate an AI generated citation in a database during the instruction session, as described in Module 05. Such a demonstration is more effective at shifting behavior than any verbal explanation, because it makes the abstract failure mode experiential. In order to make the demonstration most effective, choose a topic relevant to the course and use a general purpose AI tool the students are already using, which prevents the dismissive response that a different tool would not produce the same errors.

Third, **verification as a professional norm**: positioning citation verification not as a corrective response to AI errors but as a professional expectation that applies to all research workflows. For example, framing verification as "the step every professional researcher takes before citing" rather than "the step you need to take because AI is unreliable" produces more sustainable behavior because it connects the practice to professional identity rather than to AI specific anxiety. Such framing is accurate; professional researchers do verify citations routinely, and it positions the librarian as teaching a universal research skill rather than managing an AI problem.

Additionally, the reference consultation workflow itself can be updated to integrate AI awareness. When a patron brings a research question, a natural part of the interaction is now asking: "Have you started working with any AI tools on this project? Let's take a look at what you have." Such a question normalizes the librarian's role in AI assisted research without assuming either that the patron has or has not used AI, and it opens the opportunity to redirect problematic workflows before they produce academic integrity issues or wasted research effort. Indeed, the library's capacity to provide this kind of consultative guidance in AI assisted research is one of the strongest arguments for the profession's continued relevance in an information environment where patrons have increasing access to powerful self service tools.

Faculty conversations increasingly raise a related but distinct question: whether AI detection tools can identify whether a student's writing was AI generated. ACRL's Knowledge and Understanding competency addresses this directly, noting that AI detection tools are not completely accurate and can be circumvented. In order to advise faculty responsibly on this point, librarians should be prepared to say plainly that a detection tool's report is not evidence in the way a verified citation is evidence: false positives flag legitimate student writing as AI generated, false negatives miss AI generated text that has been lightly edited, and paraphrasing tools designed specifically to evade detection are widely available. For example, a faculty member who receives a high "AI probability" score from a detection tool and treats that score as proof of an academic integrity violation is relying on a tool with the same fundamental unreliability this module has described in AI generated content itself, just pointed in the opposite direction. The library's role in this conversation is not to validate or dismiss detection tools but to bring the same calibrated skepticism to them that this curriculum applies to AI output generally: useful as one input among several, never sufficient on its own, and never a substitute for the human judgment, a conversation with the student, an examination of their process and drafts, that an academic integrity determination actually requires.

#### **PRACTITIONER NOTE**

One structural change that proves effective is adding a source check as a standard opening step in any consultation where a patron mentions AI use. When a student says they used AI to find sources, asking to see their list and looking up one citation together takes approximately five minutes and almost always reveals at least one fabricated or inaccurate citation. That single step accomplishes more instruction than any amount of prior explanation. It also establishes, early in the consultation, that the librarian's role is not to warn the student about AI but to work alongside them as a research partner. In order to implement something similar without restructuring every consultation workflow, simply adding one opening question to the standard consultation intake makes a measurable difference: 'Have you used any AI tools so far in this project?' The answer shapes the rest of the interaction more usefully than almost any other piece of information, and it signals to patrons that AI related questions are expected and welcome in the library, not treated as a confession.

#### **KEY TAKEAWAYS**

- AI is most reliable in research support when working with documents the patron has already retrieved; summarizing provided content carries far lower hallucination risk than asking AI to find sources.
- AI cannot search databases: when asked to find peer reviewed articles, it generates seemingly credible citations from training data, many of which will not correspond to real publications.
- Search strategy development is the highest value research support application; use AI to generate term clusters, Boolean combinations, and subject heading suggestions, then take those to the actual database.
- AI native tools (Elicit, Connected Papers, Consensus) ground responses in actual scholarly databases and carry substantially lower hallucination risk than general purpose AI for bibliographic work.
- Platform embedded AI assistants, including Primo Research Assistant, Summon Research Assistant, JSTOR, EBSCO AI Insights, Ebook Central, and Scopus, ground responses in the library's own subscribed content; Scopus's agentic Deep Research feature shows where this is heading by displaying its multistep reasoning on screen, though the final synthesis still requires the same evaluation as any other AI output.
- AI detection tools are not completely accurate and can be circumvented, per ACRL's Knowledge and Understanding competency; treat detection scores as one input for academic integrity conversations, never as proof on their own.
- A workflow map (AI for topic development and search strategy, databases for retrieval, librarian judgment for source evaluation) is the most practical instructional tool for communicating appropriate AI use in the research process.
- The librarian's instruction role has expanded: teaching research AI literacy through workflow instruction, live citation demonstration, and verification as a professional norm is now a core component of single shot sessions and reference consultations.

## References

### APA 7TH EDITION

1. Association of College and Research Libraries. (2025, October). *AI competencies for academic library workers*. American Library Association. <https://www.ala.org/acrl/standards/ai>

## ACRL AI Competencies covered

Use & Application

Analysis & Evaluation

Sub-competencies: 4.1, 4.4, 3.2 · ACRL AI Competencies (2025)