

Module 02

Talking to AI effectively

Most people underuse AI because they talk to it like a search engine. This module teaches you how to actually communicate with AI, the skill that separates people who get useful results from those who don't.

By Yulia Brusova

25 min read · Reviewed June 2026

WHAT YOU'LL BE ABLE TO DO

- 1 Write a prompt that includes role, task, context, and format instructions
- 2 Use follow up prompts to refine and improve AI output rather than starting over
- 3 Set up a system prompt or custom instructions for a recurring library task
- 4 Apply a chain of thought prompt to a library evaluation or decision making task
- 5 Identify at least three patron privacy risks in AI prompting and apply strategies to mitigate them

The most significant skill gap among librarians beginning to use AI is not about knowing which tool to select. It is about knowing how to communicate with it. A vague prompt produces a vague answer; a specific, context rich prompt produces something one can actually use. In order to close that gap, this module examines the structure of effective prompting, the prompt patterns that serve recurring library tasks reliably, and the practices, including patron privacy, that distinguish responsible AI use from careless use.

The anatomy of a good prompt

A useful prompt has four elements. Not every task requires all four, but understanding them allows a librarian to construct prompts deliberately rather than by trial and error.

Role: Tell the AI who it is. For example: "You are an experienced academic librarian working at a community college with a large first generation student population." Role instructions orient the model toward the vocabulary, concerns, and register appropriate to the professional context.

Task: Tell it exactly what you want: "Write a 200-word email to faculty explaining our new database access policy." The task statement should be specific enough that there is only one reasonable interpretation of what is being asked.

Context: Give it the information it needs: "The email is for a first year writing course whose instructor has not responded to previous outreach and may be skeptical about library instruction." Context is where most librarians underinvest. The more relevant background the model has, the less generic the output will be.

Format: Tell it how to structure the output: "Use a friendly but professional tone. Three short paragraphs. No bullet points." Without format instructions, AI defaults to whatever it has learned is typical for the task type, which may not match what the librarian actually needs.

Consider, for example, the difference between these two prompts: a vague request such as "Write me an email about our library databases" and a specific one such as "You are a reference librarian at a community college. Write a friendly 150-word email to first year students introducing them to three library databases they will use for English Composition papers: JSTOR, Academic Search Complete, and ProQuest. Include one sentence about how to get help. No jargon." The second prompt produces something one can send with minimal editing. Such specificity is the core skill this module develops, and it applies equally to every library task type.

A framework for description: product, process, and performance

Dakan and Feller organize communication with AI into three layers that are worth internalizing as a framework (Dakan & Feller, *AI Fluency: Framework & Foundations*, Anthropic Academy, 2025). The first is **product description**: what, exactly, should the AI produce? This means stating not just the topic but the format, the intended audience, and the purpose. For example, asking for a research guide aimed at adult returning students who have never used a library database gives the AI something workable; asking for "a library guide" does not. The second layer is **process description**: how should the AI approach the task? This might mean specifying that the AI should consider three possible structures before choosing one, or flag any place where it is uncertain rather than fill the gap silently. Such instructions change the quality of the reasoning behind the output, not just the output itself. The third layer is **performance description**: how should the AI behave throughout the exchange? This is the layer most often skipped. It means specifying whether the AI should challenge the framing or follow the lead, whether responses should be brief or detailed, and whether the reasoning should be shown or only the result. Taking a moment to define all three layers before beginning a complex task costs one minute and reliably saves considerably more.

Iteration is the skill

Most users send one prompt, receive one response, and either accept it or abandon the attempt. This approach does not reflect how AI works most effectively. Iteration, the act of following up, refining, and redirecting, is where the genuine value lies, and it is the practice that most librarians do not develop without deliberate effort.

After receiving a first response, consider follow up prompts such as: "Make this shorter: two paragraphs instead of four"; "This sounds too formal: make it warmer and more approachable"; "The second section does not address what I need; here is what I actually require: [specifics]"; or "Give me three alternative versions of just the opening sentence." Additionally, one might ask the AI to reframe the same content for a different audience entirely, which often produces a more useful draft than revising the original.

In order to use iteration effectively, one must continue the existing conversation rather than beginning a new one. The AI retains context within a session, and that accumulated context shapes the quality of subsequent responses. For example, if earlier in a conversation the librarian has established that the audience is returning adult students who are skeptical of library instruction, follow up prompts benefit from that established context without the librarian needing to repeat it. Such continuity is a significant advantage that users who restart conversations repeatedly do not access.

Prompt patterns for common library task types

The most efficient way to develop prompting skill is to work with proven prompt structures that can be adapted to recurring library tasks rather than constructed from scratch each time. For example, a reference librarian who has developed a reliable prompt pattern for drafting research guides does not need to think through the prompt structure each time; she adapts the pattern to the specific topic and proceeds. Such patterns represent accumulated professional knowledge about how to communicate effectively with AI for specific task types.

Several patterns serve library practice reliably:

The draft-and-refine pattern (for drafting emails, instruction content, policy text): "You are a [role]. Write a [format] for [audience] that [purpose]. Tone: [description]. Length: [target]. Include: [specific elements]. Exclude: [things to omit]." This is the workhorse pattern for most library communication tasks.

The summarize-and-extract pattern (for processing long documents): "Here is a [document type]. Summarize it in [number] sentences for [audience]. Then extract: (1) the key policy changes, (2) any deadlines, (3) any action required from library staff." For example, this pattern is especially useful when working with accreditation reports, vendor contracts, or administrative memos that a librarian needs to act on quickly.

The explain-to-audience pattern (for instruction and patron communication): "Explain [concept] to a [audience description] who [relevant context about their knowledge level]. Use plain language. No jargon. Give one concrete example. Limit to [word count]." Such a prompt

reliably produces patron facing explanations that are genuinely accessible rather than merely simplified.

The generate-options pattern (for brainstorming): "Generate [number] different [options] for [goal]. Each option should be distinct from the others. Format as a numbered list with a single sentence rationale for each." For example, this pattern works well for generating five different research question framings for a patron who is stuck, or ten different workshop titles for a programming series.

In order to build a personal prompt library, the recommended approach is to document patterns that produce reliable results and save them in a shared document accessible to the whole library team. Such documentation serves both as an efficiency tool and as an institutional knowledge base that persists beyond any individual librarian's tenure, and it provides a practical foundation for the prompt library covered in Module 10.

Chain of thought prompting for complex decisions

A simple but powerful technique that most librarians do not discover without being told: asking the AI to reason through a problem step by step before producing a conclusion consistently improves output quality on complex tasks. For example, instead of asking "Should we purchase a subscription to this new database?", a more effective prompt is: "I am evaluating a new database subscription for our academic library. Here is the vendor's pitch and pricing: [paste]. Think through the following considerations step by step before giving a recommendation: (1) alignment with our subject areas, (2) cost relative to comparable resources, (3) data privacy terms, (4) ILL implications, (5) likely faculty and student usage." The output produced by this prompt is more structured, more complete, and considerably easier to act on than a response to the shorter version.

Such prompting, sometimes called chain of thought prompting, works because it constrains the model to examine each dimension before arriving at a conclusion. A model asked simply for a recommendation will produce a response that sounds confident but may skip important considerations. A model asked to reason step by step is required to work through the problem visibly, which allows the librarian to evaluate whether the AI has understood the question correctly before relying on the output. Additionally, the reasoning itself is documented, which supports professional accountability when decisions need to be explained to administrators or colleagues.

In order to apply this technique to library practice, adding the phrase "think through this step by step" or "reason through the following considerations before responding" to any prompt

that involves a decision, an evaluation, or a multistep task consistently produces more structured, more auditable output than a direct question does. Indeed, for any decision that will be shared with a supervisor or committee, the step by step reasoning is often more valuable than the conclusion itself.

System prompts and custom instructions

Most AI tools allow users to set persistent instructions, specifically the text that applies to every conversation automatically. This is where a librarian tells the AI what it should always know about the context of the work, eliminating the need to re-explain context in every session.

For library practice, useful custom instructions include information such as institution type and student population (for example, "I work at a community college serving many first generation college students and students returning to higher education after a gap"), professional role ("I am a reference and instruction librarian responsible for both walk in reference and embedded instruction"), preferred output style ("Always use plain language; avoid jargon; use active voice"), and explicit exclusions ("Never suggest citing Wikipedia as a primary source in a research context; always recommend library databases for academic sources").

Setting this up once eliminates the need to re-explain context in every conversation. Such persistent instructions are especially valuable for librarians who use AI daily for recurring task types, and they represent one of the most significant efficiency gains available at no additional cost. In ChatGPT, this feature appears as "Custom Instructions" in account settings. In Claude, it functions as "Custom Instructions" or as a Project with persistent context attached. In Gemini, it is accessible through account settings as well. Furthermore, Claude's Projects feature allows the librarian to attach documents such as a collection development policy, an instruction framework, and a list of available databases that persist across all conversations within the project. Indeed, this single setup step often produces a measurable improvement in output quality across all subsequent uses.

When to give AI information versus ask for information

There are two fundamentally different modes of AI prompting, and understanding which mode a given task requires changes how one writes the prompt and evaluates the output.

Information in: The librarian shares a document, email, or draft and asks AI to work with it. For example: "Here is a LibGuide I drafted. Improve the clarity of the introduction and suggest

more descriptive section headings." This mode is considerably more reliable because AI is working with content the user has provided rather than generating facts from training data. Such grounding in provided materials reduces hallucination risk substantially, because the model cannot fabricate details that are directly in front of it.

Information out: The librarian asks AI to produce information it does not already have in front of it. For example: "What are the most useful databases for nursing research?" This mode requires more verification because AI is drawing on training data that may be outdated, incomplete, or imprecise regarding specialized library resources. The model may confidently name databases that have changed, merged, or been discontinued since its training cutoff.

In order to get the most reliable results in daily library workflows including drafting, editing, summarizing, and brainstorming, one should default to the information in mode. This means pasting in the document one is working on, sharing a draft rather than asking for one from scratch, and providing the patron's question verbatim rather than paraphrasing it. Such an approach produces better output and requires substantially less verification. Moreover, it reflects a sound professional principle: AI is most useful as a collaborator working with what the librarian brings to the conversation, not as an independent source of facts.

Common prompting mistakes and how to avoid them

There is no doubt that certain prompting mistakes are nearly universal among librarians beginning to use AI, and identifying them directly is more efficient than allowing practitioners to discover them through accumulated frustration.

Prompts that are too short. The most frequent mistake is submitting a two or three word query, the kind of thing one might type into Google. For example, "email about databases" will produce a generic response that requires extensive revision. A prompt specific enough to produce useful output is almost always several sentences long, with role, task, context, and format specified. Such short prompts underutilize AI's capacity to work with detailed professional context.

Accepting the first response. Most users treat the first AI response as the output. In practice, the first response is a starting point. For example, after receiving a draft instruction email, effective users follow up: "The second paragraph is too formal: make it warmer"; "Shorten this by half and keep only the most essential information"; "Give me an alternative opening that leads with the patron's need rather than the library's services." Such iteration is where the quality gain happens.

Asking for too many things at once. A prompt that asks AI to "write an instruction email, design a lesson plan, and suggest five assessment activities" will produce mediocre results across all three. It is more effective to ask for one thing, evaluate and refine the result, and then proceed to the next task. Additionally, breaking complex tasks into steps allows the librarian to verify each stage before proceeding to the next.

Explaining again context in every new conversation. Users who begin a new chat session for each task lose all accumulated context. Setting up system prompts or custom instructions, as described above, eliminates this inefficiency. Such persistent context is the single most underused feature among librarians who use AI regularly, and correcting this habit alone often produces a measurable improvement in daily output quality.

Beyond these foundational mistakes, several more specific failure patterns appear regularly enough to name directly. When a response is too generic, the cause is almost always insufficient context about the specific situation. Adding details about the audience, the constraints, or the purpose makes the task genuinely specific, and a response to a genuinely specific prompt is never generic. When a response is the wrong length, specifying it explicitly in the prompt produces more consistent results than asking AI to adjust afterward: "two paragraphs" or "under 100 words" gives the model a concrete target, whereas "make it shorter" leaves the definition of shorter entirely to the model. When the tone is wrong, the most effective correction is to provide an example of writing in the tone the librarian actually wants, either as a pasted sample or as a specific description such as "the way a colleague would explain this at the reference desk." Such concrete guidance produces better calibration than adjectives like "more conversational" alone.

Several additional practical techniques from the AI Fluency course apply directly to library work (Dakan & Feller, Anthropic Academy, 2025). The first is providing context about who is asking and why, not just what is needed. For example, a prompt that identifies the librarian as preparing a ten-minute database orientation for students who have never used academic sources produces a more specifically useful response than one that simply names the topic. The second is showing rather than describing when a particular output style matters - pasting in an example of the institutional voice and asking the AI to match it is more reliable than attempting to describe that tone in the abstract. This is especially relevant when drafting patron-facing text that needs to fit an existing library communication style. Additionally, breaking a complex task into sequential steps and stating those steps explicitly in the prompt reduces the AI's tendency to collapse or skip stages. And when an interaction is not producing what is needed, the most efficient move is often to ask the AI to help improve the prompt itself - describing the actual goal and asking for a more effective way to phrase the request. This technique is underused and consistently valuable.

Evaluating AI for your specific workflows

An important practice that most librarians do not encounter without deliberate instruction is developing systematic intuition for how AI performs on the specific tasks that constitute one's actual professional work. AI is not uniformly capable or incapable across all tasks. A model that produces excellent reference email drafts may need considerable guidance for subject specific instruction content, and a tool that summarizes accreditation reports reliably may struggle with specialized library policy questions. The only way to know this for a given workflow is to test it directly.

A practical approach begins with gathering a small set of examples: five to ten instances of a task one does regularly, whether reference consultation summaries, instruction planning notes, LibGuide introductions, or patron communication drafts. Writing prompts designed to generate similar outputs and comparing the results to one's own examples reveals where AI adds genuine value, where it requires significant editing to reach a usable state, and where professional review is essential before any output can be used. For example, a cataloging librarian who tests AI generated subject heading suggestions against records she has already cataloged will develop concrete, firsthand intuition for accuracy rates in her specific collection areas. Such firsthand knowledge is considerably more reliable than general claims about AI performance for any given task type.

Such testing does not require technical infrastructure. What it requires is the discipline to compare AI output against one's own professional judgment on tasks one knows well. The result is not a pass or fail verdict on the tool but a calibrated understanding of where AI assistance is worth incorporating into a workflow and where the verification investment would eliminate the efficiency gain. In order to develop this calibration systematically, a useful practice is to identify one recurring task and run ten examples through the AI tool, noting where the output required significant correction and where it did not. Such a practice converts an abstract question about AI reliability into specific, actionable knowledge about one's actual work.

Patron privacy and what not to share in a prompt

Libraries have a long and principled commitment to patron privacy, and that commitment must extend to AI use. This is not a theoretical concern. For example, if a reference librarian pastes a patron's verbatim question, along with identifying details about the patron's research, into a commercial AI tool, that information may be processed by the AI company's servers, logged, and potentially used for training or reviewed by company employees, depending on the tool's terms of service and privacy settings. Such a practice would

constitute a meaningful breach of patron privacy in most library contexts, particularly in jurisdictions where patron records carry legal protection.

In order to use AI ethically in reference and patron services, librarians should establish clear practices about what information may and may not go into a prompt. As a general rule, remove all identifying information before pasting patron questions into AI. For example, instead of submitting "A nursing student named Maria asked about finding articles on postpartum depression for her capstone project at St. Louis Community College," the prompt should read: "A nursing student is looking for peer reviewed articles on postpartum depression for a capstone project. Suggest three search strategies and two relevant databases." Such anonymization preserves the utility of the AI assistance while protecting the patron.

Furthermore, librarians should understand the privacy terms of any AI tool used for professional work. Free tiers of commercial AI tools typically reserve more rights over user inputs than paid institutional plans do. ChatGPT's Team and Enterprise plans, Claude's Team and Enterprise plans, and Google Workspace editions of Gemini all offer stronger data protection terms than their free consumer counterparts, including commitments not to use inputs for training. Indeed, the privacy conversation between librarians and their institutions about which AI tools are approved for professional use is as important as any prompting skill, and librarians are well positioned to lead it given their existing expertise in data governance and patron confidentiality.

PRACTITIONER NOTE

A reliable prompt practice for reference work involving unclear patron questions: when a patron's research need is underspecified, remove any identifying details and paste the question into Claude with the following prompt: 'A patron has asked the following research question. Generate five clarifying questions I could ask to better understand their actual need: [question].' This approach surfaces angles that would not have been considered independently, particularly for research topics outside one's subject expertise. Such a practice takes approximately thirty seconds and consistently improves the quality of the reference interaction that follows, which is a reasonable return on a very small investment of time.

KEY TAKEAWAYS

- Effective prompts include four elements: role, task, context, and format. Specificity is the core skill.
- Iteration through follow up prompts produces far better results than accepting a first response and starting over.
- System prompts and custom instructions eliminate the need to restate your professional context in every session.
- Default to the information in mode: give AI your documents to work with rather than asking it to generate facts from scratch.
- Patron privacy obligations apply to prompts: strip all identifying information before including any patron related content.
- Chain of thought prompting, which asks AI to reason step by step, consistently improves output on complex or multipart tasks.

References

APA 7TH EDITION

1. Dakan, R., & Feller, A. (2025). *AI fluency: Framework & foundations* [Online course]. Anthropic Academy. <https://anthropic.skilljar.com/ai-fluency-framework-foundations>

ACRL AI Competencies covered

Knowledge & Understanding

Use & Application

Sub-competencies: 4.3, 2.1 · ACRL AI Competencies (2025)