

Level 2: Applied

All Librarians

Module 13

AI, labor & the library worker

Every AI adoption decision is also a labor decision. This module covers ALA's Labor value: why AI must not justify job cuts or deskilling, how to build genuine worker input into AI decisions, the hidden human labor behind the tools, and how to turn efficiency gains into better working conditions rather than fewer positions.

By Yulia Brusova

22 min read · Reviewed July 2026

WHAT YOU'LL BE ABLE TO DO

- 1 Frame AI adoption as a labor decision, not only a technology decision, using the ALA Labor value
- 2 Distinguish legitimate AI assistance from AI used to justify staff reductions or deskilling
- 3 Build meaningful worker input into AI decisions through advisory groups and consultation
- 4 Ask vendors the labor-ethics questions that surface the human work behind an AI product
- 5 Recognize patron data as extracted labor and account for it in adoption decisions
- 6 Convert verified efficiency gains into better working conditions rather than fewer positions

Most conversations about AI in libraries are framed as technology decisions, but nearly every one of them is also a labor decision, and treating it as only the former is how libraries end up in trouble. When a library adopts an AI tool, it is deciding what work will be done by whom, whose expertise will be relied upon and whose will be quietly set aside, and whether the efficiency the tool produces will be returned to the workers as better conditions or extracted from them as fewer positions. The American Library Association's guidance on artificial intelligence names Labor as one of its six core values precisely because these questions do not answer themselves, and because the history of automation in other sectors shows that efficiency gains flow to workers only when someone deliberately arranges for them to. This module is written for the library worker who wants AI adoption to strengthen the profession rather than hollow it out, and it treats labor not as an obstacle to progress but as the frame that makes progress worth having. The ALA guidance states plainly that this value 'upholds human agency over artificial intelligence and automation,' and everything that follows is an attempt to make that abstract commitment concrete in the daily decisions of an academic library.

Why labor is an AI question: human agency over automation

The reason labor belongs at the center of any library AI conversation is that automation does not distribute its consequences evenly, and the people who absorb the cost are rarely the

people who make the decision. When a vendor demonstrates a tool that drafts catalog records or answers reference questions, the efficiency is presented as a pure gain, but every efficiency in a staffed institution is also a question about what happens to the hours the tool frees and to the person who used to fill them. The ALA guidance is explicit that its Labor value 'upholds human agency over artificial intelligence and automation,' which means that the human capacity to decide, to override, and to remain the author of the work is not a courtesy the technology grants but a principle the institution must protect. For example, a cataloging assistant that suggests subject headings is compatible with human agency only when the cataloger retains full authority to accept, edit, or reject each suggestion, and it undermines human agency the moment the workflow is configured so that accepting the suggestion is the path of least resistance and questioning it is the exception. Such a distinction is invisible in a vendor demonstration and decisive in daily practice.

The specific danger the guidance identifies is that automation 'shifts decision-making to the workflow end,' which it describes as 'extremely detrimental, especially in areas, such as cataloging, where careful decisions must be made.' This is a precise and underappreciated point. When a human catalogs a record from the beginning, the careful decisions happen at the start, informed by the cataloger's knowledge of the collection and its users. When AI produces the record first and a human reviews it at the end, the careful decision has been relocated to a moment when the reviewer is checking a finished artifact rather than composing one, and the cognitive posture of checking is systematically weaker than the posture of deciding. For example, a cataloger reviewing a hundred AI-generated records at the end of a workflow is far more likely to approve a plausible error than a cataloger who would never have made that error while composing, because reviewing invites acquiescence in a way that authorship does not. Such a shift is not a failure of individual diligence; it is a structural property of where the decision sits in the workflow.

This is why the labor question cannot be reduced to job counts alone, though job counts matter. The deeper issue is agency: whether the library worker remains the author of the professional judgment or becomes the reviewer of a machine's output, and whether the institution has arranged its workflows to preserve the former or to drift toward the latter. In my practice, the tools that have strengthened my work are the ones I direct and the ones whose output I compose with, while the arrangements I am most wary of are the ones that would reposition me from the person who makes the catalog decision to the person who signs off on it. The professional judgment about which arrangement a given tool creates, and whether that arrangement preserves the human agency the work requires, belongs to the workers who do the work, and it is exactly the judgment this value exists to protect.

The displacement question: AI must not justify job cuts or deskilling

The most direct labor harm the ALA guidance addresses is the use of AI 'to justify restructuring, job reductions, deskilling, and workload shifts,' and it is worth stating the library profession's position on this without hedging: the guidance opposes it. The concern is not hypothetical. In sectors where automation has arrived, the efficiency narrative has repeatedly been used to convert a tool that could have improved working conditions into a rationale for reducing headcount, and libraries operating under chronic budget pressure are exactly the institutions where an administrator might reach for AI as a reason to leave a vacant position unfilled. The guidance is direct that libraries should not 'replace reference, readers' advisory, instructional services, or community support with AI systems,' and that 'tasks requiring empathy, judgment, and subject-specific knowledge should not be automated.' For example, a library that responds to a retirement by deploying an AI chatbot in place of a reference librarian has not achieved an efficiency; it has degraded a service that depends on the human capacities the guidance names, and it has done so in a way that is difficult to reverse once the position is gone.

Deskilling is the quieter and in some ways more corrosive version of the same harm, because it does not eliminate the position but hollows out the expertise the position was built on. When AI performs the parts of the work that developed a professional's judgment, and the human is left to review output rather than to exercise the skill, the skill atrophies, and the profession loses not a job but a competence. The guidance is emphatic that libraries must 'preserve core expertise including cataloging, subject knowledge, reference, and community support,' and it frames this preservation as an active obligation rather than a passive hope. For example, a cataloging department that routes all original cataloging through an AI tool and reserves humans for review will, over a few years, find that its catalogers have lost the fluency that made their review valuable in the first place, which is the mechanism by which deskilling becomes self-fulfilling. Such an erosion is slow enough to be invisible in any single quarter and total enough to be irreversible over a career.

The guidance also insists on a point that reframes the entire efficiency conversation: that when efficiency gains are real and verified, they should be used 'for improved working conditions and services, not staff reductions.' This is the crux. An AI tool that genuinely saves a reference librarian thirty minutes a day has created a choice, not a mandate, and the profession's position is that the recovered time should be returned to the work that only humans can do, the deeper research consultation, the instruction session, the collection development that budget pressure has crowded out, rather than converted into a justification

for asking fewer people to do more. In my practice, the honest test of any tool is whether it lets me do the parts of my job that matter more, and the honest test of any administration is whether it treats the time a tool saves as capacity to be reinvested in the profession or as slack to be cut. That judgment, about whether AI is being used to strengthen the work or to justify diminishing the people who do it, is one that library workers must be positioned to make and to name.

Worker voice: staff input, advisory groups, and bargaining units

A recurring failure in institutional AI adoption is that the people whose work will change are the last to be consulted, and the ALA guidance treats this as a defect to be corrected rather than an ordinary feature of organizational life. It calls for 'meaningful staff input into decisions about AI systems affecting their work' and recommends the use of 'staff advisory groups or governance processes to review proposed tools.' The word meaningful is doing real work in that sentence, because there is a common substitute for genuine consultation, the informational meeting held after the decision is made, that satisfies the form of worker involvement while denying its substance. For example, an administration that announces a selected AI cataloging tool and then invites catalogers to a training session has not obtained staff input; it has scheduled compliance, and the catalogers who could have identified the tool's failure modes before purchase were never asked. Such a sequence is the norm precisely because it is easier, and the guidance's insistence on meaningful input is an argument for the harder and better path.

The practical form of meaningful input is a governance structure that involves affected workers before adoption, not after. A staff advisory group with a real mandate to review proposed AI tools, to test them against actual workflows, and to recommend adoption, modification, or rejection is the mechanism the guidance points toward, and it works because the people who do the work know things about the work that no vendor demonstration and no administrative summary will surface. For example, a reference librarian on such a group will know that a proposed AI tool fails on exactly the kind of messy, half-formed question that students actually ask, a limitation that is invisible to anyone evaluating the tool on clean queries, and that knowledge is available to the institution only if the librarian is in the room before the contract is signed. Additionally, the guidance calls on libraries to 'identify work remaining human-led versus AI-assisted,' which is itself a decision that should be made with the workers whose expertise defines where that line belongs, rather than imposed by a party who does not do the work.

Where a library's workers are represented by a bargaining unit, the guidance goes further and states that when AI 'materially affects employment,' libraries must 'consult affected workers and bargaining units,' which recognizes that decisions with employment consequences belong within the structures workers have built to represent their collective interest. This is not an obstacle to good adoption; it is a safeguard that makes good adoption more likely, because a decision that can survive consultation with the people it affects is a more durable and better-reasoned decision than one that could only be implemented by avoiding them. In my practice, the AI conversations that have gone well are the ones that started with the workers and moved outward, and the ones I have watched go badly are the ones that started with a purchase and worked backward to the people. The judgment about which tools serve the work belongs, first and most legitimately, to the people who do it, and building the structures that make their voice meaningful is the institutional expression of the human agency this value protects.

The hidden labor behind the tools: data-labeling, moderation, and ghost work

The AI tools a library evaluates arrive looking finished and autonomous, but every one of them rests on a foundation of human labor that the product is designed to make invisible, and the ALA guidance asks libraries to look at that foundation before purchasing. Large language models are trained and refined by enormous amounts of human work, the labeling of data, the ranking of outputs, and the moderation of the most disturbing content the model must learn to avoid, and this work is frequently performed under conditions that would trouble any institution committed to labor ethics. For example, reporting by Time magazine in January 2023 documented that workers in Kenya were paid on the order of two dollars an hour to label graphic and traumatic content in order to make a widely used AI system safer, work that scholars including Mary Gray and Siddharth Suri have described as 'ghost work' because it is essential to the product and erased from its presentation. Such labor is the literal substrate of the tools libraries are being sold, and the guidance's inclusion of it is a refusal to pretend the tools appear from nowhere.

The guidance translates this concern into a concrete procurement practice. Before purchasing, it asks libraries to 'request documentation about data-labeling and content-moderation labor conditions,' to 'request documentation of worker protections, compensation, and mental health support,' and to 'request supply chain information including training data origins,' and it states that libraries should 'avoid tools when vendors cannot provide sufficient ethical labor documentation.' This gives the abstract concern a workable form: the same procurement conversation in which a library asks about privacy and pricing

can and should ask who labeled the training data, under what conditions, and with what protections, and a vendor's inability or unwillingness to answer is itself informative. For example, a vendor who can speak in detail about model performance but cannot describe the labor conditions behind the training data is revealing a supply chain it has chosen not to examine, and the guidance treats that absence as a reason for caution rather than a neutral gap.

I want to be honest that this is the labor dimension where library leverage is most limited, because a single academic library is a small customer to a large AI vendor and cannot, on its own, reform a global supply chain of data labor. The guidance acknowledges this directly when it notes that 'individual libraries cannot generate systemic pressure alone' and commits ALA to 'direct advocacy with vendors' and coalition work. This does not make the individual library's question pointless; it makes it part of a larger accumulation of pressure, and it aligns library purchasing with the profession's stated values even where a single purchase cannot change the vendor. For example, when many libraries ask the same labor-ethics questions in procurement, the questions become a market signal that a vendor eventually has to answer, which is the mechanism by which distributed institutional pressure becomes systemic change. In my practice, asking the question even when I cannot compel an answer has value, because it establishes that these conditions are visible to the profession and refuses the invisibility the product was designed to maintain. The judgment about whether a tool's human supply chain meets the profession's ethical standard is one libraries should insist on making, even where making it is difficult, because the alternative is to benefit from labor conditions the profession would never accept in its own building.

Patron data as uncompensated labor

There is a form of labor in the AI economy that libraries are positioned not only to witness but to enable, and the ALA guidance names it with unusual directness: the patron contributions used to train and improve AI systems are, in its words, 'often uncompensated, uncredited, and extracted as free labor.' This reframing is worth sitting with, because it identifies as labor something that is usually described in the softer language of data. When a patron's searches, questions, reading choices, and interactions with a library system are fed into a vendor's model to make that model more capable, the patron has performed work, the generation of the human behavior the model learns from, and that work has been taken without payment, credit, or in most cases awareness. For example, a discovery system that improves its recommendations by learning from every patron's search behavior is being trained by the unpaid labor of the library's community, and the value that labor creates accrues to the

vendor rather than to the patrons who produced it. Such an arrangement is easy to overlook precisely because the extraction is frictionless and the contribution is invisible.

The library's role in this is not incidental, because the library is frequently the institution that collects the patron behavior and the party that decides whether a vendor may use it. This is what makes the labor framing actionable rather than merely descriptive. The guidance's insistence, in its no-harm principle, that patron data must not be 'used for training without consent' is the library's point of leverage, because the library administers the relationship in which that consent is or is not obtained. For example, a library that reviews a discovery vendor's contract and learns that patron queries feed the vendor's model training has discovered a labor extraction it is positioned to refuse, by requiring that the feature be disabled, that the data be excluded from training, or that patrons be given a genuine opportunity to opt in rather than a default that opts them in silently. Such a refusal is the library exercising a responsibility that the labor framing makes visible: it is not only protecting privacy in the abstract but declining to conscript its community into unpaid work for a commercial model.

The connection to the profession's older commitments is close, because the labor framing and the privacy framing point to the same protective action from different angles. A library that prohibits patron reading histories and reference interactions from entering AI training systems is simultaneously protecting patron privacy and refusing to extract patron labor, and the two obligations reinforce each other. In my practice, thinking of patron data as labor rather than only as information has sharpened how I read a vendor contract, because it turns an abstract data-flow clause into a concrete question about whose work is being taken and who benefits. The professional judgment about whether the library will permit its community's behavior to be extracted as unpaid training labor is one the library actually holds, through the contracts it signs and the defaults it accepts, and exercising that judgment on the community's behalf is a direct expression of the stewardship the profession owes the people it serves.

Doing right by workers in practice: paid learning time, retraining, and redirected gains

The labor value would be an empty gesture if it stopped at prohibition, and the ALA guidance does not, because it specifies the affirmative obligations a library takes on when it adopts AI in ways that affect the people who work there. The first is that libraries must 'provide paid time for workers to learn AI functions and ethics,' which recognizes a reality that unfunded expectations obscure: learning to use AI tools well, and to understand their ethical

dimensions, is work, and requiring staff to acquire that competence on their own time is itself a labor harm dressed up as professional development. For example, an administration that expects reference staff to become proficient with an AI tool but schedules no paid time for them to learn it has quietly transferred the cost of the adoption onto the workers, who must either absorb it as unpaid effort or fall short of an expectation the institution created. Such an arrangement is common, and the guidance's insistence on paid learning time is a direct correction to it.

The second affirmative obligation concerns transitions, because AI adoption sometimes genuinely changes what a job consists of, and the guidance requires that libraries 'establish retraining and transition support before substantially changing job tasks.' The word before is essential. Retraining offered after a role has been hollowed out is a severance formality; retraining offered before the change, as preparation for a redefined role the worker will still hold, is the institution honoring its obligation to the people whose work it is reshaping. For example, a technical services department introducing AI-assisted metadata workflows honors this obligation when it retrain its staff into the higher-judgment reviewing and exception-handling roles the new workflow requires, and violates it when it changes the work first and addresses the human consequences afterward. Such sequencing is the practical difference between adoption that carries workers forward and adoption that leaves them behind, and it is a difference the institution controls entirely.

The organizing principle beneath both obligations is the one the guidance returns to repeatedly: that verified efficiency gains should be used 'for improved working conditions and services, not staff reductions,' and that libraries should 'assess labor impacts before adoption or expansion.' This is where the module's argument becomes a practice a library can actually adopt, because it prescribes a sequence: assess the labor impact before adopting, involve the affected workers in the assessment, fund the learning the adoption requires, retrain before roles change, and treat any efficiency the tool produces as capacity to reinvest in the work and the workers rather than as a reason to reduce them. For example, a library that runs this sequence around an AI cataloging tool ends up with catalogers who were consulted, trained on paid time, moved into roles that use their judgment more fully, and freed to do the collection work that mattered but never had hours, which is the version of AI adoption the profession should want. In my practice, the presence or absence of this sequence is the clearest signal of whether an institution is using AI to strengthen its people or to diminish them. The judgment about which of those an adoption represents is one library workers are entitled to make and to insist upon, because human agency over automation, the principle this entire value defends, means nothing if the humans it protects are not the ones deciding how the automation is used.

PRACTITIONER NOTE

The moment this became concrete for me was a budget season when a vacant position in our department went unfilled and the quiet suggestion in the room was that new AI tools would 'help absorb' the gap. No one said the tool was replacing the position, because no one had to; the sequence spoke for itself, and the recovered capacity that could have been reinvested in the work was instead being counted as slack that justified not hiring. What changed the conversation was not an argument about AI at all, but a labor question asked out loud: if the tool saves us time, is that time going back into the reference and instruction work we have been too short-staffed to do, or is it becoming the reason we stay short-staffed. Naming it that way made visible a decision that had been drifting toward the worse outcome by default, and it is the question I now bring to every AI adoption conversation, because the technology never answers it and the people affected by it usually are not asked.

KEY TAKEAWAYS

- Nearly every library AI decision is also a labor decision about whose work is done, whose expertise is relied on, and whether efficiency is returned to workers or extracted from them; ALA names Labor a core value because these questions do not answer themselves.
- AI must not be used to justify job cuts or deskilling; reference, readers' advisory, instruction, and community support should not be replaced, and core expertise like cataloging and subject knowledge must be actively preserved, since automation that shifts decisions to the workflow end weakens the very judgment it relies on.
- Meaningful worker input means consultation before adoption, not training after the decision; staff advisory groups with a real mandate, and bargaining-unit consultation when employment is materially affected, surface failure modes and knowledge no vendor demonstration provides.
- Every AI tool rests on hidden human labor, the data-labeling and content moderation documented as underpaid 'ghost work'; procurement should request documentation of labor conditions, worker protections, and training-data origins, and treat a vendor's inability to provide it as a reason for caution.
- Patron searches, questions, and behavior used to train vendor models are uncompensated, uncredited labor extracted from the community; the library

administers the consent relationship and can refuse to let its patrons be conscripted into unpaid training work.

- Doing right by workers is a sequence: assess labor impact before adopting, involve affected workers, fund paid learning time, retrain before roles change, and reinvest verified efficiency gains in the work and the people rather than in staff reductions.

References

APA 7TH EDITION

1. American Library Association. (2026). *Guidance on the use of artificial intelligence in libraries*. <https://www.ala.org/tools/standards-and-guidelines/guidance-use-artificial-intelligence-libraries>
2. Association of College and Research Libraries. (2025, October). *AI competencies for academic library workers*. American Library Association. <https://www.ala.org/acrl/standards/ai>
3. Association of Research Libraries. (2024). *ARL guiding principles for artificial intelligence*. <https://www.arl.org/resources/arl-guiding-principles-for-artificial-intelligence/>
4. Perrigo, B. (2023, January 18). OpenAI used Kenyan workers on less than \$2 per hour to make ChatGPT less toxic. *Time*. <https://time.com/6247678/openai-chatgpt-kenya-workers/>
5. Gray, M. L., & Suri, S. (2019). *Ghost work: How to stop Silicon Valley from building a new global underclass*. Houghton Mifflin Harcourt.
6. Williams, A., Miceli, M., & Gebru, T. (2022, October 13). The exploited labor behind artificial intelligence. *Noema Magazine*. <https://www.noemamag.com/the-exploited-labor-behind-artificial-intelligence/>
7. Clarivate. (2025). *Pulse of the library 2025*. <https://clarivate.com/pulse-of-the-library/>

ACRL AI Competencies covered

Ethical Considerations

Use & Application

Sub-competencies: 1.1, 1.4, 1.5, 2.4 · ACRL AI Competencies (2025)