



April 2026

Co-Designing a Justice-Centered Organizational AI Use Policy

Co-Designing a Justice-Centered Organizational AI Use Policy

Our fourth session in the Equity & AI Learning Series focused on how organizations can make intentional, values-aligned decisions about AI use. Facilitated by Dr. Nasreen Rajani (she/her) (<https://nasreenrajani.com/>), the session explored how to move beyond high-level AI ethics toward practical, organization-level guidance—specifically through the development of AI use policies. A justice-centered lens grounded the discussion, and it combined a short presentation with an interactive co-design process using a Miro board. Participants shared real challenges, reflected on their current practices, and contributed ideas for what meaningful, usable AI policies should include. The focus was on moving from abstract principles to practical policy that organizations can use to guide decisions, reduce risk, and stay aligned with their values.

Why This Session Matters

Organizations are being pushed to adopt AI—often quickly, and without clear guidance. Many teams are already using AI in informal ways: drafting content, analyzing data, supporting proposals. At the same time, there are real concerns around bias, confidentiality, data use, and serious social and environmental harms. The question is no longer just “Should we use AI?” It is:

- What are we using it for?
- Where does it align—or conflict—with our values?
- Who is involved in these decisions?
- What are we not willing to use it for?

They are questions about responsibility, power, and accountability—especially for organizations working with communities most impacted by harm. They are also about how AI can be used in ways that uphold values, reduce harm, and contribute positively to more just and equitable outcomes. AI use policies are one of the most practical ways to address this gap. When done well, they:

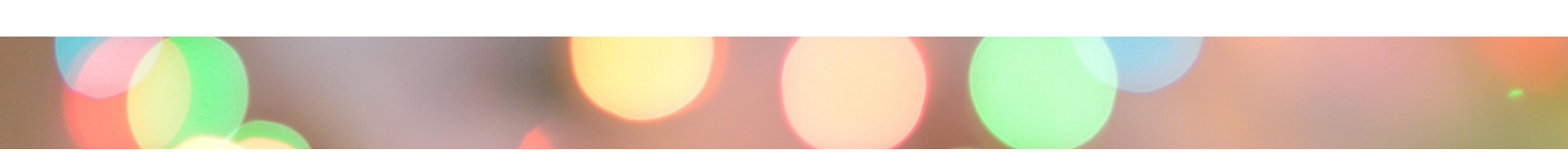
- Provide clear guidance for staff
- Create consistency across teams
- Set boundaries around risk and potential harm
- Support transparency with partners and communities
- Help ensure AI is used in ways that align with and advance organizational values

This session focused on how to create justice-centered AI use policies that are usable, grounded, and aligned with organizational values—moving beyond compliance toward more intentional, responsible, and equitable use of AI.

What Organizations Are Struggling With

The group discussion surfaced concrete challenges organizations face when trying to implement policies—not just AI policies, but policies more broadly. These insights point directly to what makes policies usable—or not. Participants highlighted several recurring challenges:

- Policies quickly become outdated as tools and use evolve
- Lack of clarity on what “responsible AI use” actually looks like in practice
- Policies are not embedded in day-to-day workflows
- Limited ownership across teams—policies are developed but not actively used
- Difficulty onboarding new staff and keeping everyone aligned
- Policies are not consistently communicated or revisited



A Practical Guide: How to Build an AI Policy in Your Organization

During the session, we explored key steps for building an AI use policy—what to include, how to approach it, and what helps make it practical and usable. The guide below is a synthesis of those discussions, drawing on insights from the group, Nasreen’s experience working on justice-centered AI policy design, and key elements from the AI Policy Toolkit for Canadian Non-Profits:

<https://mitchellconsultingsolutions.ca/resources/ai-policy-toolkit/#toolkit>

Step 1: Ground in Your Mission and Purpose (Not Tools)

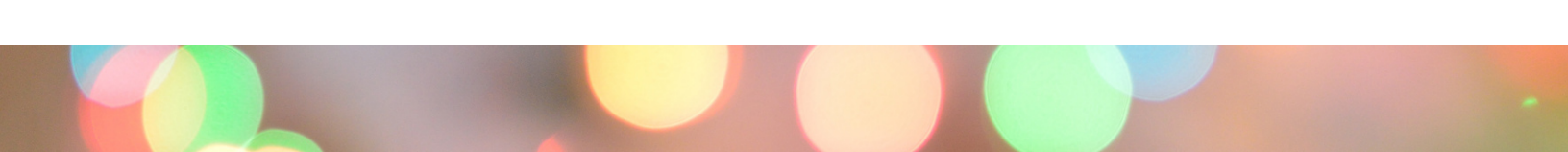
Before writing anything, start by clarifying why you are using AI and how it connects to your mission, values, and responsibilities. This includes identifying the problems you are trying to solve, understanding where AI may add value, and where it may introduce risk or tension—particularly in relation to the communities you serve. A key point raised in the session was that this process should not happen in isolation. Justice-centered approaches to AI policy design emphasize creating space for communities, civil society, and impacted groups to shape decisions based on lived experience—not just technical expertise. For example, initiatives like the People’s AI Consultation (<https://www.peoplesaiconsultation.ca>) offer structured facilitation guides, discussion formats, and guiding questions to support inclusive, participatory policy development. Grounding your approach in this way helps ensure that AI use is not only aligned with your mission, but also accountable to those most affected by it. If AI does not clearly support your work—or cannot be used in a way that reflects your values—it may not be appropriate to prioritize it.

Step 2: Define Your Core Principles

Strong AI policies are grounded in a small set of clear, guiding principles that reflect your organization’s values and help anchor decision-making. These principles act as a reference point when navigating uncertainty or trade-offs. Common principles include human-centered use (AI supports—not replaces—people), equity and fairness (actively identifying and mitigating bias), transparency (being open about when and how AI is used), privacy and data protection, and accountability (ensuring humans remain responsible for decisions and outputs). What matters most is not adopting a long list, but selecting a few principles that are meaningful in your context and can be applied in practice. You might find it helpful to draw from existing frameworks—such as the OECD AI Principles or work from the Mozilla Foundation—as starting points, and then adapt them to reflect their own mission and the realities of the communities they serve. Keeping this section concise (typically 4–6 principles) makes it easier to communicate, remember, and apply across teams.

Step 3: Map How AI Is Already Being Used

Before defining rules, it is important to understand how AI is already being used across your organization. Many teams are already experimenting—often informally—with tools for drafting content, summarizing information, analyzing data, or supporting proposals. Mapping this use helps move the conversation from abstract concerns to real decisions. A practical approach is to create a simple, shared space (such as a document or Slack thread) where staff can note how they are using AI, what is working well, and where they feel uncertain or uncomfortable.



These examples often surface important patterns—for instance, where AI is adding value, where outputs require significant correction, or where there may be risks related to data use, bias, or misrepresentation. This step is also an opportunity to identify gaps in understanding across the team and to surface differences in access or comfort with AI tools. Grounding your policy in these real experiences makes it far more relevant, and ensures that the guidance you develop reflects how people actually work—not how you assume they do.

Step 4: Set Clear Boundaries for Use (and Non-Use)

Once you understand how AI is being used, the next step is to define clear boundaries. This is where many policies fall short—relying on broad language like “use responsibly,” which leaves too much open to interpretation. Policies are most useful when they clearly outline what is allowed, what requires caution, and what is not permitted—across both use cases and tools. A practical approach is to define specific use cases and categorize them. For example, using AI to brainstorm ideas or structure content may be appropriate, while summarizing internal documents may require caution and review. Other uses—such as uploading interview transcripts, working with personal or sensitive data, or generating content that represents community voices—may be explicitly prohibited. These distinctions help staff understand not just what to do, but when and why certain uses are higher risk.

It is also important to consider boundaries related to tools themselves. Not all AI tools are equal in how they handle data, store information, or train models. Some organizations choose to limit use to a small number of approved tools, or to restrict specific features (such as uploading files or using tools that retain data). Being explicit about which tools can be used—and under what conditions—reduces uncertainty and helps manage risk. Strong policies also define areas of non-use. In some cases, the most responsible decision is not to use AI at all—particularly in contexts involving sensitive data, complex qualitative insights, or situations where accuracy, consent, and representation are critical. Naming these boundaries clearly reinforces that responsible AI use includes knowing when not to use it. The goal is not necessarily to restrict all use, but to create enough clarity that staff can make informed, consistent decisions.

Step 5: Clarify Roles and Responsibilities

One of the consistent challenges raised in the session was that policies often exist without clear ownership—leading to uneven use, uncertainty, and limited follow-through. Defining roles and responsibilities helps ensure that AI use is consistent, accountable, and supported across the organization. At a minimum, policies should outline expectations for all staff using AI tools. This typically includes using tools in line with the policy, reviewing outputs for accuracy, and exercising judgment—particularly in higher-risk contexts. It is also important to clarify the role of management in supporting implementation, including onboarding new staff, reinforcing expectations, and overseeing how AI is being used within teams.



Some organizations also identify a specific role—often referred to as a policy champion—to help keep the policy active over time. This person (or small group) may support training, gather feedback from staff, track emerging issues, and coordinate updates as tools and practices evolve. In smaller teams, these responsibilities may be shared, but making them explicit helps prevent the policy from becoming inactive or overlooked. Clarifying roles also reinforces an important principle: while AI can support work, responsibility for decisions and outputs remains with people. Making this visible in the policy helps build a culture of accountability, where staff feel supported in using AI—but also understand their role in ensuring it is used appropriately.

Step 6: Identify and Address Key Risks

A strong AI policy makes risks visible and actionable. While many risks are already understood in principle, they often remain implicit or unevenly addressed in practice. Naming them clearly helps teams recognize where extra care is needed and supports more consistent decision-making across the organization. Several key risk areas emerged in the session. These include bias and discrimination in AI outputs, particularly when working across different contexts or with underrepresented groups; data privacy and security, especially when handling personal, sensitive, or unpublished information; and the risk of misrepresentation—where AI-generated content simplifies, distorts, or speaks inaccurately on behalf of communities. Participants also noted the risk of over-reliance on AI, where outputs are accepted without sufficient review or critical thinking.

Addressing these risks does not require overly complex systems, but it does require clarity. In practice, this means linking risks to specific guidance within the policy. For example, if bias is a concern, the policy can require human review and validation of outputs. If data privacy is a risk, it can clearly prohibit entering certain types of information into AI tools. Where misrepresentation is a concern, it may be necessary to limit or avoid using AI in contexts where voice, lived experience, and nuance are critical. Importantly, risks are not the same across all organizations or contexts. Policies should reflect the specific realities of your work, including the populations you serve and the types of data you handle. Taking the time to identify and articulate these risks helps ensure that AI use is not only efficient, but also responsible and aligned with your values.

Step 7: Establish Simple Governance and Decision-Making

Once boundaries, roles, and risks are defined, the next step is to clarify how decisions about AI are made in practice. Governance does not need to be complex, but it does need to be clear. Without it, organizations often face uncertainty around which tools to use, how to respond to emerging issues, or who has the authority to make decisions. At a basic level, this means defining how new AI tools or use cases are assessed and approved. For example, some organizations introduce a simple review process before adopting a new tool—considering factors such as data handling practices, potential risks, and alignment with the policy. This does not need to be formal or time-intensive, but it should create a moment of reflection before new tools are widely used.

It is also important to establish how concerns or issues are raised and addressed. Staff should know where to go if something feels unclear, risky, or inappropriate, and feel comfortable raising questions without hesitation. This helps create a culture where AI use is discussed openly rather than handled individually or informally. Finally, governance includes how decisions are documented and revisited. Even simple documentation—such as noting why certain tools are approved or restricted—can help ensure consistency over time and support learning across the organization. As with the rest of the policy, governance should be seen as something that evolves. Starting with a simple structure and refining it as use grows is often more effective than trying to build a comprehensive system from the outset.

Step 8: Support Staff Through Training and Ongoing Guidance

Even the strongest policy will fall short if staff do not understand how to apply it in practice. Participants consistently emphasized that AI policies need to be supported by clear, ongoing guidance—not just a one-time rollout. This starts with building a shared baseline understanding of what AI tools can and cannot do, including their limitations, risks, and appropriate use cases. In practice, this means keeping training simple, relevant, and grounded in real work. Rather than focusing on theory, organizations are finding it more effective to use concrete examples—how AI might be used in a report, a proposal, or data analysis—and to walk through what is appropriate, what requires caution, and what should be avoided. Creating space for questions and discussion is equally important, as staff often encounter uncertainty in real time and need a place to explore it.

Ongoing support also matters. As tools evolve and use expands, staff will continue to encounter new situations that are not fully covered in the policy. Regular touchpoints—such as short discussions in team meetings, quick refreshers, or shared examples of AI use—help reinforce expectations and keep the policy active. Some organizations are also experimenting with lightweight approaches, such as short guides, quick reference sheets, or informal sessions, to make support more accessible. Ultimately, the goal is to build confidence and consistency. When staff understand both the intent of the policy and how to apply it in their day-to-day work, AI use becomes more thoughtful, aligned, and accountable.

Step 9: Review and Adapt Over Time

AI policies are most effective when they are treated as living documents rather than fixed rules. As tools evolve and new use cases emerge, policies need to adapt to remain relevant and useful. Participants emphasized that many policies become outdated quickly—not because they are poorly designed, but because they are not revisited. In practice, this means building in regular review points. Some organizations set a simple review cycle (for example, every 6–12 months), while others revisit the policy more informally based on emerging questions or challenges. What matters most is creating a clear expectation that the policy will be updated over time. A practical way to support this is to draw on real examples of AI use within the organization. Keeping track of how AI is being used—what is working well, where uncertainty arises, and where risks are encountered—provides a concrete basis for refining the policy. These examples help ensure that updates are grounded in actual practice rather than assumptions. Ongoing review is also an opportunity to incorporate feedback from staff and, where relevant, from partners or communities. This reinforces shared ownership and helps ensure the policy continues to reflect the realities of the work. Rather than aiming for a perfect policy from the start, organizations are finding value in starting with something simple, applying it in practice, and improving it over time. This iterative approach allows policies to evolve alongside both the technology and the organization's needs.

Making your AI Policy Work in Practice: Key Tips for Implementation

These tips are drawn directly from participant discussions and reflect what organizations are finding works in practice when trying to make AI policies usable and relevant.

- **Make It Part of Daily Work:** Build policies into existing processes and tools. Rather than expecting staff to refer back to a standalone document, integrate policy into report and proposal templates, checklists and workflows or data analysis processes. Example: Add a simple step in templates: ☐ AI used ☐ Output reviewed for accuracy ☐ No sensitive data included
- **Use Real Examples to Keep It Relevant:** Keep a running, shared list of how AI is actually being used across the team. This could be a short document or Slack thread where staff note quick examples—what they used AI for, what worked well, and what felt unclear or risky (e.g. “helpful for brainstorming report structure,” “needed heavy editing,” “not comfortable uploading interview data”).

- Reviewing these real examples every few months gives the team a practical basis to refine the policy—clarifying what to allow, what to limit, and what to avoid based on actual experience, not assumptions.
- **Be Clear and Specific:** Define specific use cases, not just general principles. Be explicit about data boundaries early on. Include a “What this looks like in practice” section directly in the policy, with clear, real examples. For instance: using AI to brainstorm or structure a report is generally appropriate; summarizing internal documents may be acceptable with caution and careful review; but uploading interview transcripts, participant names, or unpublished reports should be clearly prohibited. Rather than relying on broad statements like “protect confidentiality,” policies should spell out exactly what that means in practice. Including a few concrete scenarios for each key rule helps staff make quick, confident decisions in real situations
- **Keep It Short and Usable:** Make policies short, clear, and easy to use. Participants emphasized making policies “as seamless, clear and brief as possible.” Focus on concise language, practical guidance and removing unnecessary jargon. Example: Create a 1-page version staff can quickly reference during their work.
- **Create Shared Ownership:** Create shared ownership—not one-time consultation. Some participants noted that involving the full team—even if “a bit boring”—builds real buy-in and makes the policy feel like everyone’s responsibility.
- **Clarify Roles and Accountability:** Clear roles help ensure that policies remain active and supported over time. Defining expectations for staff, management, and any designated leads (such as a policy champion) makes responsibility visible.
- **Keep It Visible and Active:** Use simple, recurring touchpoints to keep policies active.
 - Review a small section of the policy (1–3 pages) during regular staff meetings - Example: A monthly 5-minute team prompt: “Did anyone use AI this month in a way that raised questions?”
 - Rotate responsibility across team members to lead short discussions
 - Include quick follow-ups (e.g. quiz, reflection, discussion)
 - Pair policy review sessions with snacks or interactive activities to make engagement more fun
- **Link Policy to Principles:** Policies are often seen as restrictive or irrelevant. Frame the policy as helping staff make decisions and protecting communities and data. Connect it directly to organizational values. Example: “This policy helps us ensure our work remains accurate, ethical, and grounded in community realities.”
- **Start Simple and Build Over Time:** Policies do not need to be perfect from the start. Beginning with a simple, clear approach—and refining it over time based on real use—reduces pressure and supports learning. This allows the policy to evolve alongside both the organization and the technology.

Key AI Policy Resources

The session surfaced a range of practical resources that organizations can use when developing or refining their AI policies. These are not one-size-fits-all solutions, but they offer strong starting points, examples, and language that can be adapted to your context.

1. Foundational Principles and Frameworks

These are useful when you are defining your overall approach, values, and boundaries.

- **OECD AI Principles:** A widely used international framework outlining principles for trustworthy, human-centered AI. Use this to ground your policy in recognized standards and rights-based approaches.
<https://www.oecd.org/en/topics/sub-issues/ai-principles.html>
- **Mozilla Foundation:** Offers practical, people-first perspectives on AI, including work on responsible AI and public interest technology. Use this to think critically about power, equity, and real-world impacts.
<https://www.youtube.com/playlist?list=PLnRGhgZaGeBv4FyPOmMY3OOFQJAz8J0lb>

2. Practical Organizational Policy Examples

These are helpful when you are drafting or refining your own policy.

- University of Toronto – AI Use Guidelines: Clear guidance on acceptable use, confidentiality, and staff responsibilities. Use this for concrete examples of how to structure internal guidance. <https://people.utoronto.ca/employees/administrative-staff-and-ai/guideline-on-the-use-of-artificial-intelligence-by-university-administrative-staff/>
- University of British Columbia – Principles for Generative AI Use: Focuses on responsible experimentation and risk mitigation. Use this to shape high-level principles that still allow flexibility. <https://genai.ubc.ca/guidance/principles/> [https://it-genai-2023.sites.olt.ubc.ca/files/2024/08/PRINCIPLES FOR MITIGATING RISKS.pdf](https://it-genai-2023.sites.olt.ubc.ca/files/2024/08/PRINCIPLES_FOR_MITIGATING_RISKS.pdf)
- Toronto Public Library – AI Policy: A practical, operational policy focused on staff use of AI tools. Use this as a model for simple, applicable organizational policies. <https://tpl.ca/policies-and-terms-of-use/artificial-intelligence-policy/>
- Wikimedia Foundation – AI Strategy: Includes strong language on how AI use aligns with mission and community values. Use this for inspiration on aligning AI use with purpose and public accountability. https://meta.wikimedia.org/wiki/Strategy/Multigenerational/Artificial_intelligence_for_editors
- AI Policy Toolkit for Canadian Non-Profits – Mitchell Consulting Solutions
- A comprehensive, step-by-step guide to developing an AI use policy, tailored specifically to the nonprofit sector. Covers everything from core principles and ethical considerations to governance, risk management, and implementation. Use this if you want a full structure to build your policy from scratch or strengthen an existing one. <https://mitchellconsultingsolutions.ca/resources/ai-policy-toolkit/#toolkit>

3. Applied and Sector-Specific Guidance

These resources show how organizations are integrating AI into real-world, high-stakes contexts.

- World Bank Independent Evaluation Group – AI Guidance Note: Focuses on using AI in evaluation while maintaining rigor and quality. Use this if you are working with data, research, or evaluation. https://ieg.worldbankgroup.org/sites/default/files/Data/Evaluation/files/report-artificial_intelligence_strategy.pdf
- World Bank AI Strategy (IEG): Shares how AI is being introduced across an organization, including limits and scope. Use this to think about phased or controlled adoption. <https://documents1.worldbank.org/curated/en/099136005132515321/pdf/IDU-dccd6e52-4ee3-4294-a264-28fda8a94a49.pdf>
- ICRC – Policy on Artificial Intelligence (2024): A strong example of a human-centered AI policy in complex, high-risk environments. Use this to understand how to approach AI where risks to people are significant. <https://www.icrc.org/en/publication/building-responsible-humanitarian-approach-icrcs-policy-artificial-intelligence>

Closing Reflection

AI is already being used across organizations—often informally, unevenly, and without clear guidance. What this session highlighted is the need to make intentional decisions about AI—whether to use it, how to use it, and where not to use it—based on your values and context. We hope this session—and the reflections and resources shared here—are useful as you begin or continue developing your own approach. If you are in the process of thinking through or building an AI policy, don't hesitate to reach out or continue the conversation. This is evolving work, and there is real value in learning from each other along the way.

If you have ideas for future session topics, potential speakers, relevant case examples, or questions you'd like us to explore, I would love to hear from you. Please reach out to me - Paula Richardson - at richardson@salanga.org.