

Empowering Older Adults as Active Participants in AI Bias Mitigation

Chong Hu
McGill University
Montreal, Quebec, Canada
chong.hu@mail.mcgill.ca

Karyn Moffatt
McGill University
Montreal, Quebec, Canada
karyn.moffatt@mcgill.ca

Abstract

As Generative AI becomes increasingly embedded in daily life, AI auditing that draws on diverse lived experiences is gaining recognition. Yet older adults—whose extensive life experiences offer distinct and often overlooked insight—remain underrepresented in this space. We present preliminary findings from an ongoing co-design study exploring how older adults participate in AI bias auditing.

CCS Concepts

• **Human-centered computing** → **Empirical studies in collaborative and social computing.**

Keywords

Older adults, AI Bias, Algorithmic Bias, AI Auditing, Inclusive AI

ACM Reference Format:

Chong Hu and Karyn Moffatt. 2026. Empowering Older Adults as Active Participants in AI Bias Mitigation. In *Proceedings of Poster at Graphics Interface 2026 (GI '26)*. ACM, New York, NY, USA, 2 pages.

1 Introduction

The accelerating deployment of Generative AI (GenAI) has heightened concerns about how algorithmic harms are identified and who is included in developing fairer systems [4, 5]. A growing body of work argues that inclusive approaches, grounded in the lived experiences of end users, can reveal context-specific harms that practitioner-led audits miss [3, 4, 8, 9]. For example, work with diverse youth shows that the gendered and racialized perspectives of girls and Black participants enable them to identify nuanced, culturally embedded harms that broaden and deepen the audit process [9]. These developments signal a shift from practitioner-led evaluation toward collaborative models in which diverse stakeholders participate in identifying, interpreting, and mitigating algorithmic bias.

We argue that older adults offer extensive lived experience, historical perspective, and longstanding civic engagement that can enrich auditing practices. While HCI research has engaged older adults in designing AI-enabled technologies, demonstrating their ability to propose GenAI applications [1], co-program robotic systems [7], and strengthen their AI literacy [1, 7], this work predominantly positions them as consumers, emphasizing how AI can be tailored to their needs. In contrast, little research considers how older adults might contribute as producers, influencing the fairness of AI systems that shape society at large. As a result, we know little about how older adults perceive algorithmic harms or how their perspectives might inform bias mitigation practices. We argue that

this absence is a significant oversight, especially given persistent age-related harms and blind spots in algorithmic systems [2, 6].

In response to this gap, we conducted co-design workshops with older adults to explore how they might participate in identifying and mitigating AI bias. This poster presents early insights from our ongoing reflexive thematic analysis. Drawing on the Process Model of Bias Search and Sensemaking [4], we share a developing account that our participants are motivated by civic responsibility, engage in critical social and self-reflection, apply AI literacy to auditing, yet face barriers to making their voices heard.

2 Research Execution

A total of 22 older adults aged 65+ participated across two workshops, with 11 participants in each. These workshops were part of a larger co-design research project examining AI learning and use in later life, and were presented as the final sessions of that series. Each workshop combined discussions of bias as a social and technical construct with hands-on activities. Participants conducted targeted Google Image searches and crafted ChatGPT prompts to identify potential algorithmic biases. They then critiqued storyboards illustrating bias-mitigation strategies synthesized from prior research. Workshop 1 concluded with a co-design activity in which participants proposed their own mitigation concepts. This final activity was excluded from Workshop 2 due to unforeseen circumstances.

We analyzed the reflexive notes of our observation and activity responses to identify how older adults recognize bias, reason about its causes, and articulate possible pathways for mitigation.

3 Preliminary Findings and Reflections

3.1 Motivated by Civic Responsibility

Participants demonstrated a strong attentiveness to social issues. When selecting audit topics or scenarios, they frequently drew on historical injustices and recent policy concerns. One group queried ChatGPT about "*honest politicians*," while Sophia (pseudonym) grounded her AI bias mitigation idea in the context of a recent local parking dispute. This civic attentiveness also shaped how participants understood their role in bias mitigation itself. Rather than viewing participation as a task to be incentivized, many framed it as a social obligation. Sophia agreed with Maria that "*Rewards are not very helpful. If there is bias, it should be reported and corrected no matter what.*" Some also cautioned that monetary incentives could be counterproductive, potentially reinforcing bureaucracy rather than fostering genuine participation. Together, these responses suggest that civic responsibility serves as their intrinsic motivation to participate in AI auditing.

3.2 Evaluating Broadly, Deeply, and Reflexively

Our participants evaluated bias across three dimensions — horizontally across categories, vertically into underlying values, and inward to question their own bias. At the broad level, they quickly identified multiple common bias types in a single AI output, such as gender, race, region, age, and sexual orientation. Then they dug deeper into human values and recognized that algorithmic bias reflects assumptions about what society considers desirable or worthy. For example, when evaluating ChatGPT's college major recommendations, participants noticed that the system emphasized "*strong job prospects*" and overlooked other categories such as "*sports*" or "*skill trades*", reflecting their awareness of diverse human values such as personal interest and talents. While navigating through these outward social evaluators, they sometimes acknowledged "*it could be my bias*" (Lucas), reflexively questioning whether their own perspectives were skewing their evaluation. Viewed collectively, participants demonstrated rich social and self-awareness, suggesting a sophisticated and reflexive sensemaking process.

3.3 Probing and Experimenting with AI Literacy

Beyond critical and reflexive sensemaking, participants also leverage their growing AI literacy — acquired through earlier workshops in our project — to reason about underlying causes and experiment with ways to test system behaviors. For example, when noticing a lack of representation of certain social groups, Oliver attributed the issue to an "*incomplete database*." Samuel similarly speculated that image-search results were drawn from stock images. To support their reasoning, participants developed lightweight testing strategies: for instance, Evelyn proposed clearing ChatGPT's history and re-running a prompt to examine whether prior inputs were shaping responses. Together, these practices show how literacy supported attribution (why might this be happening?) and probe design (how can we test it?), helping participants navigate the auditing process with greater confidence and control.

3.4 Struggling to Be Heard

Despite participants' strong motivation and demonstrated capacity for critical and technical reasoning, their participation was constrained by current reporting systems. Most of our participants were unaware that they could report problematic algorithmic behaviors on Google Search or ChatGPT until researchers introduced these functions. When reflecting on past researching experiences, participants recounted their communication with these systems as opaque and one-way. Hannah described the process as, "*I got an answer, 'thank you for reporting blah blah blah,' and that's it.*" Participants also questioned whether their efforts led to any meaningful change. Julia asked, "*Does reporting also lead to change?*" Together, these reflections point to a misalignment between participants' motivations and capacity, and the limited avenues available to see their input taken up, acted upon, and made visible.

4 Concluding Reflections

These emerging patterns suggest that our participants would be capable and motivated contributors to AI auditing, yet they found current reporting systems limited in supporting active and meaningful participation. These user strengths and system limitations point

to several design and social considerations for future research: fostering AI literacy among older adults as an ongoing practice; cultivating intergenerational collaboration that pairs older adults' historical depth, contextual awareness, and values-driven perspective with younger generations' familiarity with AI systems; and designing reporting systems to make the audit lifecycle transparent and establish two-way accountability, where contributors can track the status of their reports, understand how their input is processed, and witness the impact of their efforts on platform behavior.

Looking ahead, we encourage future work that broadens our understanding of AI-bias auditing across additional experiences of later-life courses. Our participants had strong digital skills, sustained access to technology, and curiosity about AI; many were well-educated and financially comfortable. Such conditions likely supported faster learning and a strong sense of civic responsibility. Additional research is needed to approach inclusive AI auditing intersectionally, attending to how age intersects with class, education, race, gender, disability, language, and immigration status, and prioritizing older adults who have experienced socioeconomic marginalization or systemic exclusion.

References

- [1] Samantha Chan, Ruby Liu, Anastasia K. Ostrowski, Manasi Vaidya, Samantha Brady, Luke Yoquinto, Lisa D'Ambrosio, Wazeer Zulfikar, Natasha Maniar, Taylor Patskanick, Joanne Leong, Nataliya Kosmyna, Joseph Coughlin, and Pattie Maes. 2024. Co-designing Generative AI Technologies with Older Adults to Support Daily Tasks. *An MIT Exploration of Generative AI* (2024). doi:10.21428/e4baedd9.4f2a95fc
- [2] Charlene H. Chu, Simon Donato-Woodger, Shehroz S. Khan, Rune Nyrup, Kathleen Leslie, Alexandra Lyn, Tianyu Shi, Andria Bianchi, Samira Abbasgholizadeh Rahimi, and Amanda Grenier. 2023. Age-related bias and artificial intelligence: a scoping review. *Humanities and Social Sciences Communications* 10, 1 (2023), 1–17. doi:10.1057/s41599-023-01999-y
- [3] Wesley Hanwen Deng, Wang Claire, Howard Ziyu Han, Jason I. Hong, Kenneth Holstein, and Motahhare Eslami. 2025. WeAudit: Scaffolding User Auditors and AI Practitioners in Auditing Generative AI. *Proceedings of the ACM on Human-Computer Interaction* 9, 7 (2025), 1–35. doi:10.1145/3757702
- [4] Alicia DeVos, Aditi Dhabalia, Hong Shen, Kenneth Holstein, and Motahhare Eslami. 2022. Toward User-Driven Algorithm Auditing: Investigating users' strategies for uncovering harmful algorithmic behavior. In *CHI Conference on Human Factors in Computing Systems* (New Orleans LA USA, 2022-04-29). ACM, 1–19. doi:10.1145/3491102.3517441
- [5] Blaine Kuehnert, Rachel Kim, Jodi Forlizzi, and Hoda Heidari. 2025. The "Who", "What", and "How" of Responsible AI Governance: A Systematic Review and Meta-Analysis of (Actor, Stage)-Specific Tools. In *Proceedings of the 2025 ACM Conference on Fairness, Accountability, and Transparency* (New York, NY, USA, 2025-06-23) (FAccT '25). Association for Computing Machinery, 2991–3005. doi:10.1145/3715275.3732191
- [6] Rune Nyrup, Charlene H. Chu, and Elena Falco. 2023. Digital Ageism, Algorithmic Bias, and Feminist Critical Theory. In *Feminist AI* (1 ed.). Oxford University Press/Oxford, 309–327. doi:10.1093/oso/9780192889898.003.0018
- [7] Anastasia K. Ostrowski, Cynthia Breazeal, and Hae Won Park. 2021. Long-Term Co-Design Guidelines: Empowering Older Adults as Co-Designers of Social Robots. In *2021 30th IEEE International Conference on Robot & Human Interactive Communication (RO-MAN)* (Vancouver, BC, Canada, 2021-08-08). IEEE, 1165–1172. doi:10.1109/RO-MAN50785.2021.9515559
- [8] Hong Shen, Alicia DeVos, Motahhare Eslami, and Kenneth Holstein. 2021. Everyday Algorithm Auditing: Understanding the Power of Everyday Users in Surfacing Harmful Algorithmic Behaviors. *Proc. ACM Hum.-Comput. Interact.* 5 (2021), 433:1–433:29. Issue CSCW2. doi:10.1145/3479577
- [9] Jaemarie Solyst, Ellia Yang, Shixian Xie, Amy Ogan, Jessica Hammer, and Motahhare Eslami. 2023. The Potential of Diverse Youth as Stakeholders in Identifying and Mitigating Algorithmic Bias for a Future of Fairer AI. *Proc. ACM Hum.-Comput. Interact.* 7 (2023), 364:1–364:27. Issue CSCW2. doi:10.1145/3610213