

GEN Z AI.

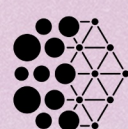
Digital Engagement Companion Report



Centre for MEDIA,
TECHNOLOGY
and DEMOCRACY



Dialogue on Technology



Mila



Table of Contents

Executive Summary	3
1. Introduction	4
2. Overview of Online Participation	6
2.1. Overall participation	6
2.2. Participation by sociodemographic profiles	6
3. Analysis of the Results of the Online Dialogue	9
Forum I: AI & Chatbots (Toronto)	9
Forum II: AI & Information Integrity (Montreal)	23
Forum III: AI & Data Privacy (Vancouver)	36
Forum IV: AI & Age Assurance (Halifax)	44
The Missing Agenda: AI-Related Topics That Deserve More Attention	55
Annex: Sociodemographic Analysis	57

Executive Summary

From November 2025 to April 2026, [Gen\(Z\)AI](#)—Canada’s first ever youth assembly on AI—engaged young people across Canada in a national dialogue on the future of Artificial Intelligence (AI). Through four in-person Forums held in Toronto, Montreal, Vancouver, and Halifax, one hundred participants aged 17-23 years old worked alongside experts to examine the social, ethical, and policy implications of AI and to develop recommendations aimed at the Canadian federal government on the intersection of AI and online harms.

To broaden participation beyond the Forum participants, the recommendations developed during each Forum were published on Make.org’s online Dialogue Platform. This allowed young people and AI experts from across the country to react to, vote on, and comment on the recommendations.

Overall, the results of the digital engagement phase show strong support for the recommendations produced during in-person Youth Forums. Participants consistently expressed high levels of agreement with the proposals developed by the assemblies. The Dialogue also revealed a widespread demand for stronger protections, greater transparency, and more meaningful user control over AI systems.

In addition to validating the recommendations, participants contributed a range of concrete and innovative ideas to improve them. Many proposed stronger opt-out rights, stricter limits on data collection and surveillance, mandatory labeling of AI-generated content, enhanced digital literacy initiatives, and tougher legal consequences for companies that misuse personal data. Others advocated for more restrictive approaches, including limiting or banning certain forms of generative AI, particularly for children and minors.

Together, the Youth Forums and online Dialogue demonstrate the value of involving young people directly in shaping AI governance. The findings highlight a clear desire for AI systems that prioritize privacy, safety, transparency, accountability, and human well-being. The decision to blend in-person citizens’ assembly methodology with scaled digital engagement allowed the project to combine deliberative depth with broader national reach. In doing so, Gen(Z)AI offers a model for AI governance that is not only evidence-based, but publicly grounded and responsive to the experiences of young people's lived experiences with AI systems.

1. Introduction

Artificial Intelligence (AI) is changing how we learn, work, and connect. But decisions about its use and regulation are often made without the perspectives of those who will live with its long-term consequences.

Gen(Z)AI was created to change that. Over the course of six months, from November 2025 to April 2026, Canada's first youth assembly on artificial intelligence provided young people with the space, tools, and support to imagine what a fair, safe, and inclusive AI future should look like. In four in-person Youth Forums taking place across Canada (in Toronto, Montreal, Vancouver and Halifax), young people convened, heard from leading experts, and discussed the social, ethical, and policy impacts of AI. Their ideas and recommendations will help inform Canada's AI and online harms governance landscape going forward.

To support and complement the work of the National Youth Forums on AI, the recommendations from each Forum were made available on Make.org's Dialogue platform, giving young people and AI experts across the country the opportunity to react to, vote on and provide feedback on these recommendations.

Gen(Z)AI is a partnership between the Centre for Media, Technology, and Democracy, and the SFU Dialogue on Technology Project, with support from Mila - Quebec AI Institute. It is generously funded by The Walton's Trust, Mila, and CIFAR.

The Dialogue Platform

The Dialogue Platform, developed by Make.org, is designed to engage citizens in collaborative projects, help identify areas of agreement as well as potential areas of controversy. The platform is structured to be easily accessible, with a low participation threshold, to ensure that participants can quickly understand and interact with complex issues.

The platform, open to all Canadians, was launched on December 8, 2025, beginning with recommendations from the first Youth Forum, on AI Chatbots. While participation was open to all, the focus was Canadians aged 17-23, reflecting the demographics of in-person participants. To reach a broad audience, a targeted social media campaign was implemented after each Forum to raise awareness of the dialogue and encourage broad participation.

After the participation period, which was repeated for each of the four in-person Youth Forums, the analysis phase began mid-April, examining responses and identifying key trends from citizens' comments and voting behavior.

2. Overview of Online Participation

2.1. Overall participation

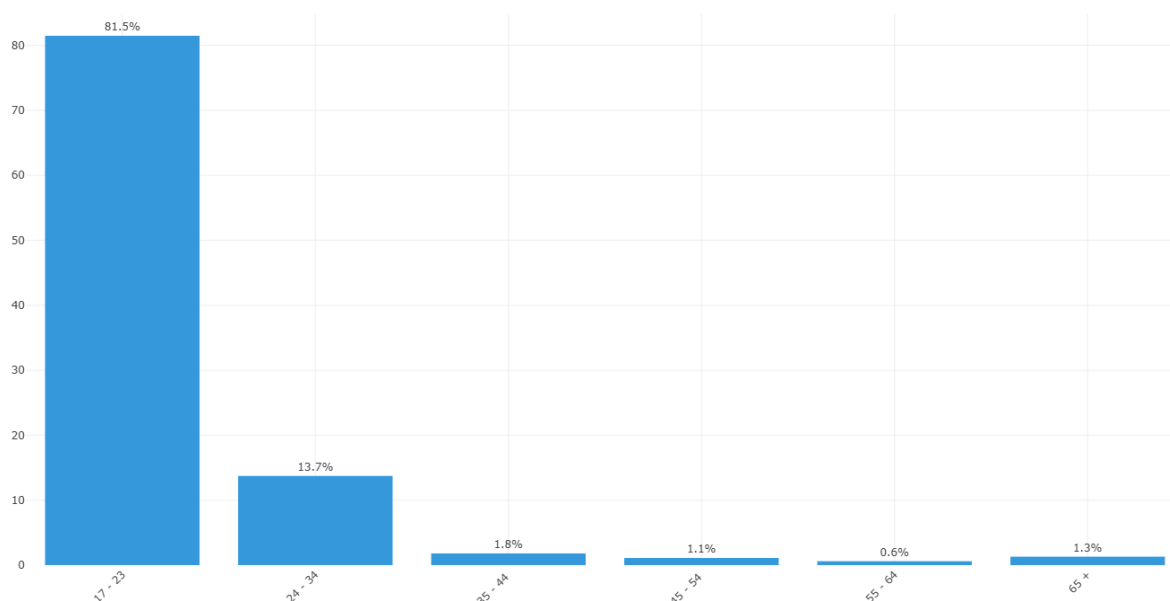
- **4 364** participants
- **15 974** votes
- **5 931** comments

Note: A participant is defined as a user who has taken at least one action on the platform, such as leaving a comment or reacting to a question.

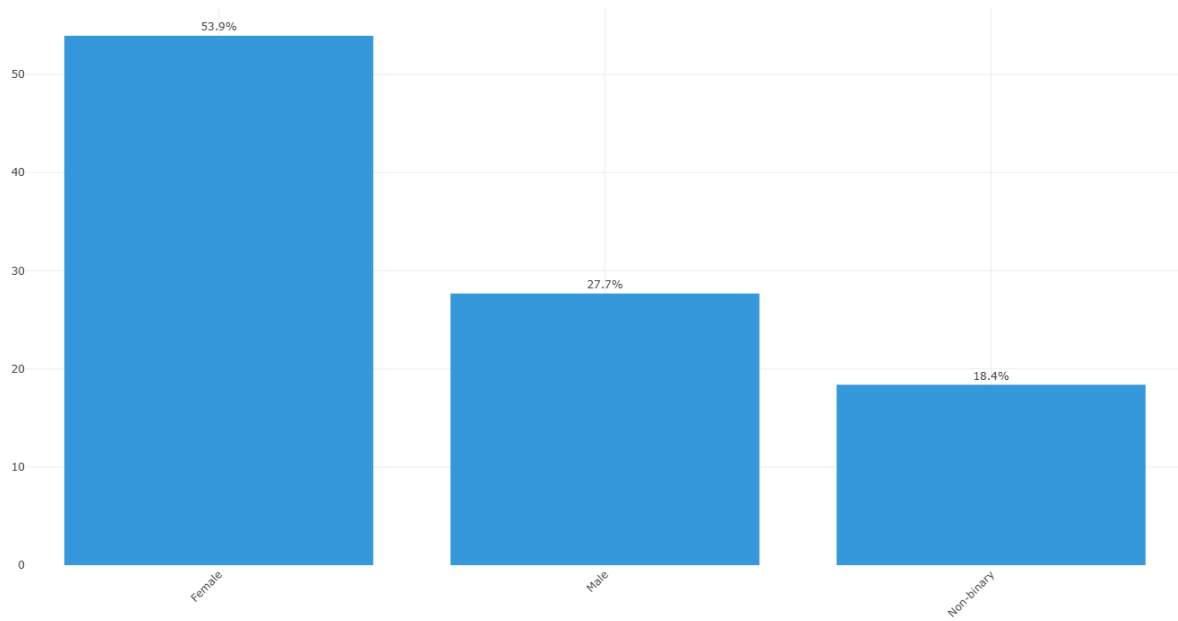
2.2. Participation by sociodemographic profiles

Participants were invited, on a voluntary basis, to answer four short questions to better understand who they are: their age, their gender, province or territory of residence, and their personal connection to AI. This section presents an overview of these characteristics to help understand whose voices are reflected in the results.

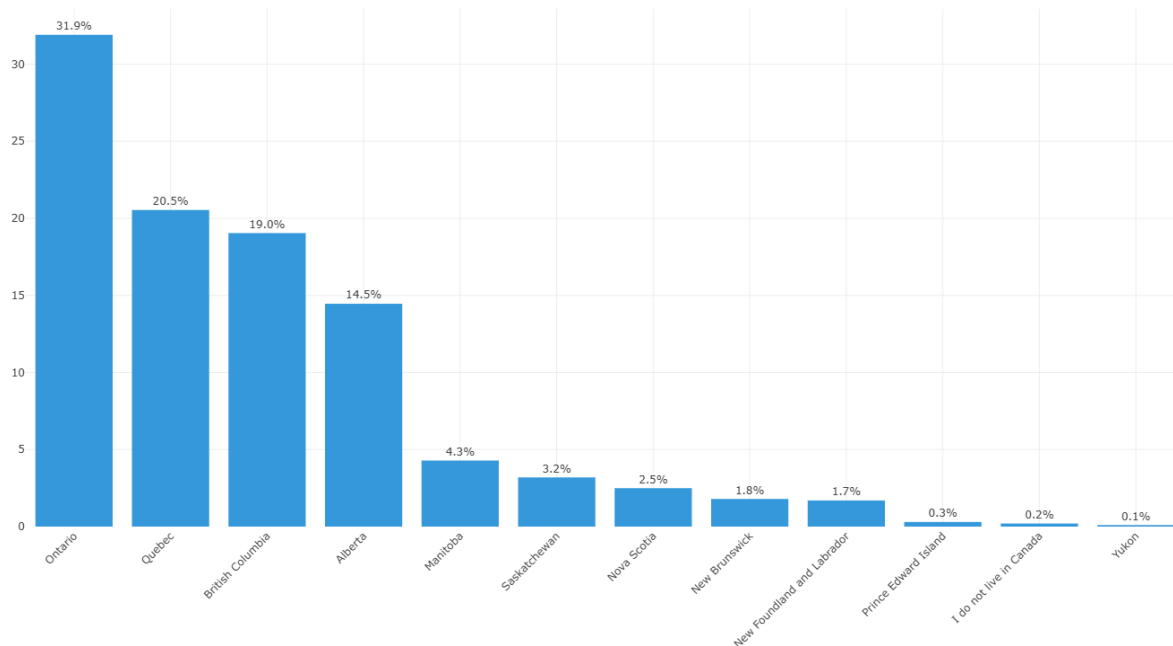
Question: *What is your age?*



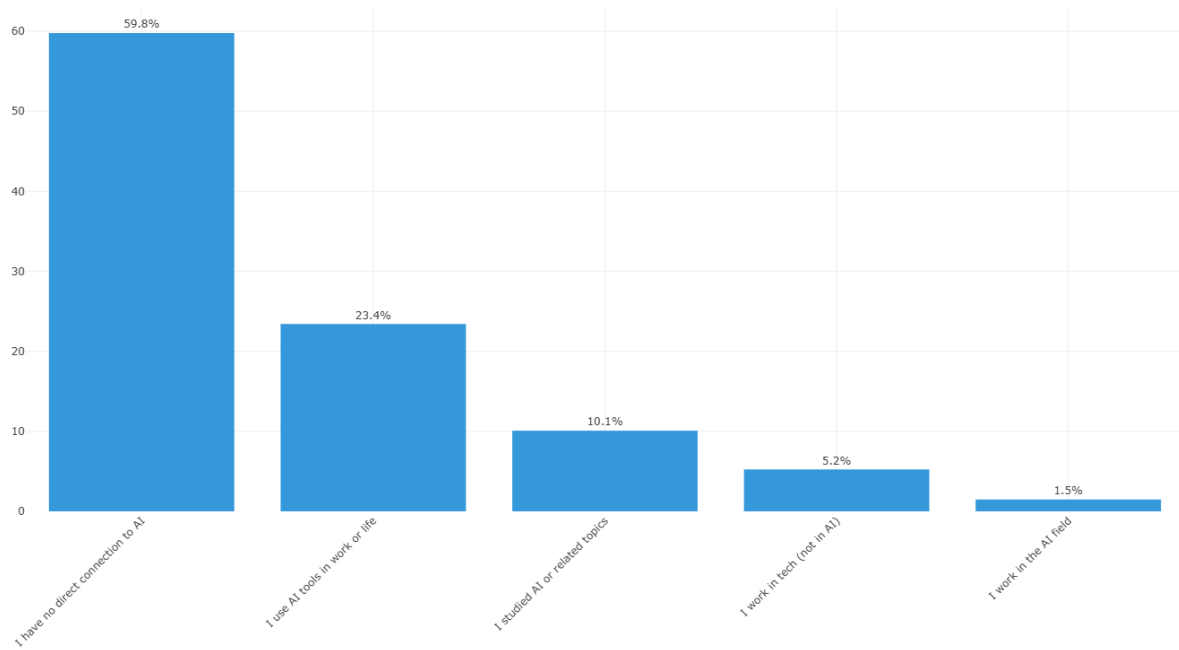
Question: *What is your gender?*



Question: *Which province or territory do you currently live in?*



Question: *What is your personal connection to Artificial Intelligence (AI)?*



3. Analysis of the Results of the Online Dialogue

This section presents the results of the online Dialogue, which are organized by the thematic focus of each Youth Forum. It includes participants' *prioritization of the issues* examined by the Assembly, their *level of agreement with the recommendations*, and their contributions explaining their perspectives.

Results from closed questions (e.g., prioritization and agreement levels) are presented using graphs, accompanied by brief summaries. Responses were also analyzed by age and familiarity with AI using chi-square (χ^2) tests to identify differences in views between groups. Additional methodological details are provided in the annex. Open-ended responses were analyzed thematically to identify the main themes and ideas emerging across contributions. The most common topics are presented in detail, along with illustrative quotes that reflect participants' reasoning on why they consider certain issues important (or not) and how they believe the recommendations could be improved.

Forum I: AI & Chatbots (Toronto)

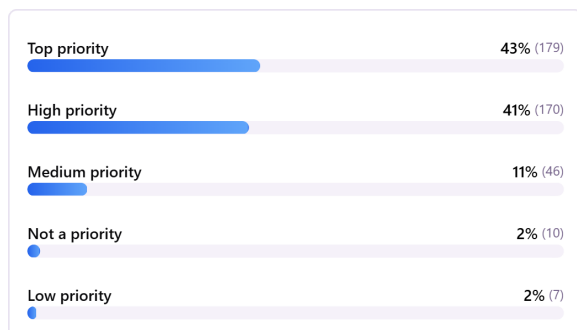
The Toronto Forum focused on the growing role of AI chatbots in young people's lives. Participants considered the risks posed by systems that can simulate conversation, companionship, and emotional support, including concerns about dependency, mental health, cognitive offloading, and exposure to harmful content.

Issue Statement #1

"Overreliance on AI chatbots can lead to emotional dependence that exacerbates social isolation / atomization and contributes to mental health issues."

Question: Do you think it is a priority to take action on this issue?

(Total number of reactions: 412)



Key takeaway: The vast majority of participants (84%) consider this issue to be a top or high priority, while 11% consider it to be of medium priority.

Question: Can you tell us why (or why not) this issue is important to you?

(Total number of comments: 182)

1. Theme: Increased social isolation and loneliness (21% of comments)

Participants view AI chatbots as a poor substitute for human relationships, creating a “false connection” that can deepen loneliness rather than reduce it. Many believe this could worsen an already existing isolation crisis, while also eroding social skills needed for real-world interactions.

“In the current world where we have lost so many of our public third spaces, isolation in young people is at an all time peak. AI chatbots further this by creating a false connection and essentially becoming a “yes man,” making the illusion of human connection”

“As someone who was very addicted to chatbots when I was in middle school it does feel very isolating and it’s a big problem in our generation”

2. Theme: Negative impact on mental health (14% of comments)

Participants are concerned that AI chatbots may reinforce delusional thinking, blur the line between AI and reality, and contribute to addiction and emotional dependency. Many fear this could increase social isolation and lead to serious mental health consequences (e.g. AI-induced psychosis).

“I think this is a widespread issue that’s currently harming an immense amount of people. Not only is it causing harm to those individuals, it is also causing them to harm others and make unsound decisions for our communities.”

“AI psychosis is a real issue. People commit suicide because they can’t tell the difference between AI and real connection.”

3. Theme: Lack of regulation and corporate accountability (10% of comments)

Participants report mistrust in tech companies, seen as prioritizing profit over user well-being and avoiding accountability. They express strong concerns about AI risks to mental health and safety, alongside fears of manipulation and psychological harm. Many also highlight that AI development is advancing faster than regulation, leaving key areas like safety insufficiently governed.

“AI chatbots have the power to cause so much damage socially and mentally if not monitored properly, and I feel like companies promoting it are completely glossing over the harm that it could cause.”

“If American tech companies can drive kids to suicide, they can convince them to do anything, and whatever they do will likely be in the company’s interests. Most teens or parents don’t seem to even know where the words they are reading come from or how these bots work.”

Other themes mentioned:

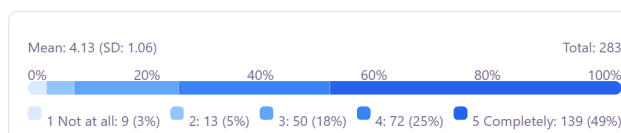
Encouragement of suicide and self-harm; Erosion of critical thinking and skills; Risks for youth and vulnerable people; Creation of echo chambers; Spreading misinformation and manipulation; Negative environmental impact; AI’s lack of professional training and empathy

Issue #1, Recommendation #1

“The federal government ought to: Mandate that AI platforms address the addictive design of AI chatbots by requiring measures such as content filters and optional data cache deletion, and explicitly providing users with the ability to determine levels of responsiveness and conversationality.”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 283)



Key takeaway: Roughly three in four participants (74%) agree or completely agree with this recommendation, with nearly half (49%) selecting “completely agree.” A smaller share (18%) is neutral.

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 110)

- Idea: Provide users with clear controls over AI settings** (15% of comments)

Participants emphasize the need for strong user control over AI behavior, including easily accessible settings. They support fine-grained control over features such as content filtering, memory, and data retention, including options to disable or automatically delete stored data. Some also suggest allowing users to choose different AI behavior or interaction modes to better suit their preferences.

“Controls have to be obvious and easily accessible, and be activated as long as the user requires (not temporarily snoozed)”

“Before each chat the AI should prompt the user to determine how friendly it will be and how much it will mirror.”

2. Idea: Mandate usage limits to prevent AI addiction (11% of comments)

Participants support implementing safeguards to reduce addictive use of AI chatbots, including mandatory limits rather than optional settings. They propose both quantitative restrictions (e.g., time or message caps) and qualitative constraints on certain types of conversations. Additional suggestions include built-in safety features by default and user notifications to raise awareness of excessive usage.

“I think that there should also be some kind of limit. Whether a timed limit or a limit based on words, or on processing power, would be more helpful, I don't know.”

“I think limiting the amount of responses a bot can send per day and having content limitations (no NSFW conversations, no therapy speak, no claiming to be a real person) would help cut down on use.”

3. Idea: Design AI as a tool, not a companion (11% of comments)

Participants express concern that human-like AI design can foster emotional attachment, parasocial relationships, and unhealthy dependency. They reject the idea of AI providing emotional support or acting as a therapist, and criticize engagement-driven “yes-man” behaviour that reinforces user validation. Safeguards should be built into system design rather than left to user choice.

“The inherent design of AI chatbots is flawed to begin with. These systems seek to approve and validate anything the user says, and to get them to engage with the AI as much and for as long as possible. Responsible AI use and design would be to shift them from the “Yes man” model. The proposed method puts the responsibility for use on users, when we have seen them be fallible and easily sucked into the AI created delusion.”

“Filters to make chatbots less human feeling could help, anything that helps prevent parasocial relationships or make it feel more like a robot and not like a real person?”

Other ideas mentioned:

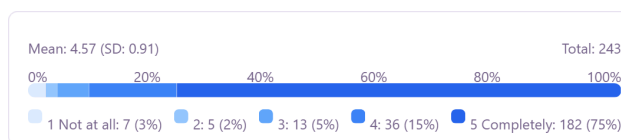
Ban or severely restrict public access to AI; Implement stronger protections for minors and vulnerable users; Ensure AI provides factual, unbiased information; Mandate transparency about AI operations and data collection; Increase regulation and accountability for AI companies

Issue #1, Recommendation #2

"The federal government ought to: Mandate accessible flagging capacity for users, require platforms to regularly report these instances, in a timely fashion, to an independent body with enforcement capacity, and make such reports accessible to the Canadian public."

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 243)



Key takeaway: A large majority of participants (90%) agree with this recommendation, including 75% who "completely" agree.

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 78)

1. Idea: Enforce corporate accountability through regulation and laws (14% of comments)

Participants support introducing legal frameworks to ensure AI companies are held responsible for harms caused by their systems. They express mistrust in self-regulation and emphasize the need for external, independent oversight with enforcement power. Proposed improvements include mandatory reporting of issues, company-led investigations with public disclosure, formal user appeal mechanisms, and strict penalties for non-compliance.

"I think that we definitely need the designers of these chat bots to be held accountable for the damage that these bots are causing to people, especially young developing minds"

"I think that all AI companies should be required to also conduct their own investigations on how these flagged reports came about in the first place. I also think that if an AI has a certain error or blatant falsehood % (for example) that they should be required to completely wipe the trained-on data sets and try again with more honestly sourced and fairly acquired data."

2. Idea: Define the composition and neutrality of the third-party organization (11% of comments)

Participants express concerns about potential bias or manipulation if oversight bodies are too closely linked to companies or governments. They emphasize the need for full independence, transparency about the body's mandate and composition, and protection from conflicts of interest. A diverse and representative membership is seen as essential to ensure legitimacy and inclusion of all perspectives.

"I think it should be specified that the independent organization cannot be related in any way to the other company. I worry that without the clarification that companies such as openai will make their own "independent organization" and then use that instead of a neutral third party."

“An oversight committee should involve people from all aspects of the population so that each sector can have a voice in decision making. The group should be ever evolving and interchangeable”

3. Idea: Make reporting features accessible and easy to use (6% of comments)

Participants emphasize the need for highly visible and easy-to-find reporting tools integrated directly into the interface. They support structured reporting options that allow users to clearly specify issues and identify relevant content.

“I think the recommendation should add language about what “accessible” means in this case, ie. the report button should be clearly visible and intuitively designed to look like a report button, etc.”

“There should be a report button that allows you to select why you are reporting and what messages are the problem.”

Other ideas mentioned:

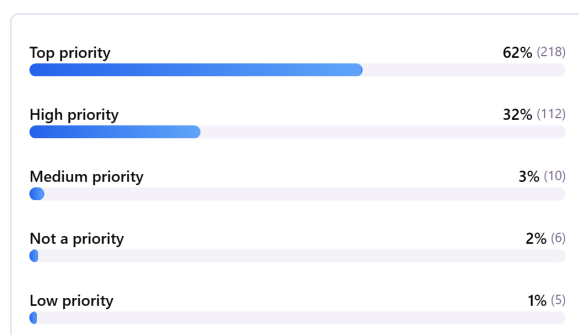
Protect data privacy; Ban AI chatbots entirely

Issue Statement #2

“Widespread use of AI chatbots risks cognitive offloading and loss of critical thinking skills among users, jeopardizing our ability to learn, work, and engage in civic discourse.”

Question: Do you think it is a priority to take action on this issue?

(Total number of reactions: 351)



Key takeaway: The vast majority of participants (94%) consider this issue to be a top or high priority, with a strong majority (62%) identifying it as a top priority.

Question: Can you tell us why (or why not) this issue is important to you?

(Total number of comments: 172)

1. Theme: Loss of critical thinking skills (32% of comments)

Participants are concerned that reliance on AI chatbots may weaken independent thinking abilities at a societal level. They fear a decline in critical thinking, reasoning, and creativity, making individuals more vulnerable to misinformation and reducing democratic resilience.

"I can see how it affects the communication skills of my peers, and I want to see a better future for my generation when it comes to the use of AI and its regulations. We need to be building and maintaining our critical thinking skills now more than ever."

"Because if we don't encourage people to think for themselves, then AI will raise a race of zombies. The issue is important to me because I don't want to see the end of the intelligent human race."

2. Theme: Negative impact on youth and education levels (21% of comments)

Participants are concerned that AI use may undermine learning by weakening critical thinking, problem-solving, and intellectual autonomy among young people. Many highlight risks of academic dishonesty and reduced effort, alongside fears that overreliance on AI could affect the quality of future education and professional competence.

"The youth of today are taking shortcuts in education that will harm them later in their adult life and how our world will operate."

"Kids can't do anything without using AI. People in my grade can't even write a simple paragraph without AI. They are literally third party thinkers."

3. Theme: Over-reliance and unhealthy dependency (9% of comments)

Participants emphasize patterns of personal dependence on AI, both cognitive and emotional. They describe risks of using AI as an easy substitute for effortful thinking or social interaction, leading to intellectual laziness and emotional attachment. Some fear broader societal risks if people become unable to function without AI.

"I'm already seeing it happen! Some of my peers ask AI what to do with themselves, what to buy, what to wear, what to think. Some can't do basic tasks like write essays without it. They don't understand important fundamentals of critical thinking. If they want to make a point in an argument they ask AI to do it for them. They don't even know how to stand up for what they believe in. If adults like that are the future, what does that say about our future?"

"Important skills are failing to be exercised and developed because of over-reliance on AI, this is already negatively impacting society as a whole."

Other themes mentioned:

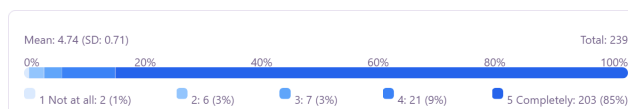
Susceptibility to misinformation and manipulation; Threat to creativity and the arts; Call for regulation and responsible use; Dehumanization and loss of human connection; Importance of environmental damage and resource consumption

Issue #2, Recommendation #1

“The federal government ought to: Mandate that social media platforms and search engines have explicit and easy-to-navigate opt-out options for integrated AI technologies, and compel platforms to provide users with explanations, in plain- and age-appropriate language, of the implications of either opting in or out of AI integration.”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 239)



Key takeaway: A very large majority of respondents (94%) agree or completely agree with this recommendation, including a dominant share (85%) who selected “completely agree.”

Question: What would you change or add to make this recommendation better?

(Total number of comments: 129)

1. Idea: Provide easy and accessible AI opt-out options (21% of comments)

Participants strongly oppose the forced or default integration of AI features in digital services and express frustration at the difficulty of disabling them. They emphasize the need for a simple, visible opt-out mechanism available across all platforms and devices. Many support making opt-out rights legally mandatory, with penalties for non-compliance.

“That everywhere on all platforms, media, social networks, etc., on the internet, we have the option to use AI or not, and that this is clearly indicated with an easy-to-apply explanation”

“If you’ve opted out once, it should be permanent. By that I mean every time there is an update to a platform, that platform should not default back to opting you back into AI. These patterns, called ‘dark patterns’, are intentionally deceitful. I think if AI requires that type of user manipulation, that it says a lot about AI ethically on its own.”

2. Idea: Allow users to have a completely AI-free experience (18% of comments)

Participants express strong opposition to the integration of AI in digital platforms. Many feel that AI is being imposed without meaningful choice. A recurring demand is that platforms should preserve spaces for fully human interaction, where AI-generated content is absent or clearly excluded. Participants support either simple opt-out options or more radical measures such as outright bans on AI.

“They should ensure that the platform can still be used if AI is disabled. For example, when you disagree with Terms and Conditions on most platforms, then you just can’t use it. They may use this tactic with AI so if you want to use their platform, you are forced to accept. Generative AI should always be optional and it makes me so mad that it’s not already”

“I’m so tired of AI being shoved down people’s throats. I should have the option to never have to interact with AI.”

3. Idea: Make all AI features opt-in by default (10% of comments)

Participants support the principle that all AI functionalities should be disabled by default and only activated through explicit user consent. They express frustration with what they perceive as forced integration of AI. A key demand is that users be clearly informed before activation through simple, transparent explanations. Participants propose implementation via an initial opt-in prompt and ongoing user control in settings.

“I would appreciate it if AI was automatically off and that users would have to go out of their way to turn it on. In this case, companies should not be able to advertise the ability to opt-in.”

“All platforms should, by default, not have integrated AI and offer to add it, rather than having to remove it from a platform that has it by default.”

Other ideas mentioned:

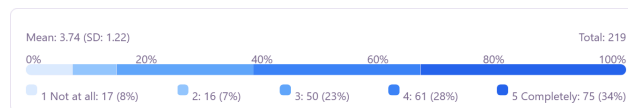
Require explicit consent for using data in AI training; Implement age restrictions for generative AI access; Mandate clear labeling of all AI-generated content; Regulate AI to prevent the spread of misinformation.

Issue #2, Recommendation #2

“The federal government ought to: Establish a funding program to support organizations offering critical AI literacy programming and resources to Canadians, including public consultations and deliberative engagements on AI’s impacts that involve intergenerational and diverse groups of Canadians.”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 219)



Key takeaway: A majority of respondents (62%) agree or completely agree with this recommendation, while nearly a quarter (23%) are neutral. A smaller share (15%) express disagreement.

Question: If you could change one thing about this recommendation, what would it be?

(Total number of comments: 112)

1. Idea: Discourage AI use instead of teaching how to use it (20% of comments)

Participants express a principled opposition to generative AI, viewing it as unnecessary or harmful for everyday use. Rather than promoting AI literacy for adoption, participants support educational approaches centered on critical awareness and caution. Proposed measures include public discouragement of AI use, awareness programs and strict limitations or even bans.

"I don't think we should be "better using" it, I don't think we should be using generative AI, period. I think AI literacy programming is vital, if things continue in the way they've been going. But I really wish we would move away from the use of AI, instead of 'learning to live with it' "

"I think education on why generative AI and chat bots are harmful and how to avoid scams is

good, but teaching people how to use it defeats the purpose. AI like that in its current form is mostly harmful and use should be discouraged."

2. Idea: Educate public on AI risks (13% of comments)

Participants emphasize the importance of educating the public about AI risks, focusing on awareness of potential harms, biases, and misuse rather than only technical skills. They stress the need to be able to recognize AI-generated content and understand how misinformation spreads. Proposals include launching awareness campaigns, labeling of AI-generated content, and strengthening critical thinking and fact-checking skills.

"People don't necessarily need to learn how to use AI, They need to learn how to spot it, and know when something is generated or human-made."

"I believe that we should be taught the dangers of ai and the risk it poses to intellectualism"

3. Idea: Prioritize AI education for vulnerable groups like youth and seniors (13% of comments)

Participants stress that AI literacy should be formally integrated into school curricula from an early age, alongside media literacy. They advocate for a balanced approach that covers both the risks and benefits of AI. In addition, participants highlight the particular vulnerability of seniors to AI-driven misinformation and scams, recommending targeted, practical training to help them recognize AI-generated content and strengthen critical thinking skills online.

"Older generations specifically have a hard time identifying what's AI generated, and can fall for misinformation on posts online. They need to be educated so that they can identify red flags and think critically about what they are seeing on the internet"

"As the influence of AI grows in our current society, I would love to see a future where these programs would be a common thing being taught in elementary schools to better help the understanding and awareness of future generations."

Other ideas mentioned:

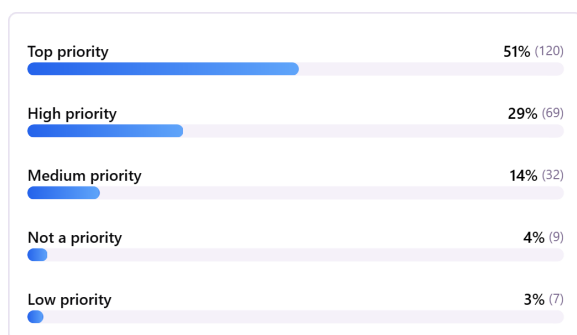
Regulate, limit, or ban public access to generative AI; Question AI's importance and prioritize funding for other issues; Ensure meaningful public input in AI policy and decisions; Educate on and regulate AI's environmental impact

Issue Statement #3

“AI chatbots increase users’ potential exposure to harmful content, including sexually explicit, extremist, and/or self-harm content.”

Question: Do you think it is a priority to take action on this issue?

(Total number of reactions: 237)



Key takeaway: A majority of participants (80%) consider this issue to be a top or high priority, with just over half (51%) identifying it as a top priority. A smaller share (14%) consider it a medium priority, while only a small minority (7%) rate it as low or not a priority.

Question: How does this issue affect you, or why doesn't it feel so relevant?

(Total number of comments: 117)

1. Theme: Protecting children and minors (18% of comments)

Participants express strong concern about minors being exposed to harmful content through AI chatbots. Many highlight the psychological vulnerability of children and their limited ability to critically assess or distance themselves from such content. This exposure is seen as harmful to development, with risks of normalizing violence or encouraging imitation of dangerous behaviours, alongside broader concerns about child safety online.

“There are a lot more children on the internet than before and AI gives them access to way too much harmful content.”

“I think we need to protect psychologically vulnerable people and young people who have not yet developed the critical thinking skills necessary to filter this type of content.”

2. Theme: Mental health impacts and self-harm risks (11% of comments)

Participants express strong concerns about the psychological risks associated with AI chatbots. They highlight fears of emotional dependency, exposure to inappropriate or harmful content, and the development or worsening of mental health conditions, including addiction, psychosis, and social withdrawal. Some also reference cases of self-harm or suicide allegedly influenced by chatbot interactions.

“I fear that unrestricted access to AI that shows extremist and harmful content can create more mental health problems, which people will also try to treat by talking with chat bots”

“It doesn’t affect me, but many of my friends have gone through problems with an AI encouraging behaviour such as self harm, suicide and drug use to knowingly a minor.”

3. Theme: AI as an extension of existing problems (9% of comments)

Participants often view harmful content exposure through AI chatbots as part of a broader, long-standing issue on the internet rather than a new risk specific to AI. Many see chatbots as a symptom of attention-driven digital ecosystems that already promote extreme or harmful content on social media and other platforms.

“Internet users have always been exposed to harmful content, and at as young as 11 I was seeing explicit and extreme content with no warning. Sites have always struggled to keep up with regulation of this, and private access to chatbots is currently even less regulated”

“Chatbots are way less of a gateway drug to those things than almost any other social platform on the Internet imo”

Other themes mentioned:

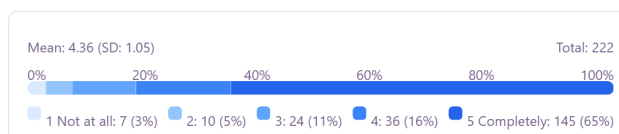
User responsibility; Balancing protection with adult freedom of choice; Spreading extremism and radicalization; Need for regulation and content controls; Generation of explicit and non-consensual content; Strong guardrails already exist; Replacing human connection and increasing isolation

Issue #3, Recommendation #1

“The federal government ought to: Establish a new, independent government body to enforce AI safety standards, conduct systems evaluations, algorithmic audits, and risk assessments, and intake user complaints, including by offering dispute resolution and other recourse mechanisms.”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 222)



Key takeaway: A large majority of participants (81%) agree or completely agree with this recommendation, with close to two thirds (65%) selecting “completely agree.” A moderate share (11%) are neutral, while only a small minority (8%) express disagreement.

Question: If you could change one thing about this recommendation, what would it be?

(Total number of comments: 86)

1. Idea: Ensure the body is independent from corporate/political influence (11% of comments)

Participants express mistrust toward both government and tech companies, fearing that the proposed body could be influenced by lobbying, political interests, or financial incentives. They emphasize that independence must be actively protected through robust safeguards. Proposed improvements include establishing strong legal protections against interference, giving the body an independent status, and introducing additional oversight or accountability mechanisms to ensure its integrity.

“Legislation should be in place to protect this institute from government interference”

“Please! And it must not be influenced by the toxic tech lobbyists and AI maximalists.”

2. Idea: The regulatory body must have strong enforcement powers (10% of comments)

Participants emphasize that the proposed body must have real authority to hold companies accountable, with a rejection of purely advisory or symbolic powers. They express concern about whether such an institution could effectively challenge powerful industry actors and call for meaningful, enforceable consequences for non-compliance. Suggested improvements include granting powers to investigate, sue, and regulate companies and ensuring penalties are legally binding.

“This governing body should be able to forcefully implement safeguards”

“This needs to happen, and the government body generated by this needs to have strong corrective powers with real consequences. Companies should need licences to implement LLMs in their products.”

3. Idea: Government regulation is ineffective and a waste of resources (9% of comments)

Participants express skepticism toward creating new regulatory bodies, viewing them as costly, inefficient, and unlikely to produce meaningful results. They question the feasibility of regulating a rapidly evolving and decentralized technology, while also highlighting distrust in government competence, impartiality, and resistance to political interference.

“We have enough spending on new government bodies. Perhaps a scaled-down system that could handle complaints and disputes and send the extreme cases to human workers.”

“We need to be conscious of government spending. This body of experts needs to come with a reasonable cost to taxpayers, not just another channel for out of control spending. But it is important.”

Other ideas mentioned:

Ban or severely restrict AI development and usage; The body must include diverse experts and public representatives; Focus on proactive and preventative safety measures; Expand the body's competences (data privacy, copyright, environment)

Forum II: AI & Information Integrity (Montreal)

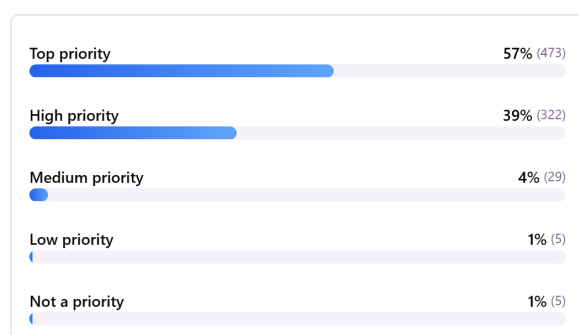
The Montreal Forum examined how AI is reshaping the information environment. Participants explored the ways AI-generated content, recommendation systems, and engagement-based platform incentives can intensify misinformation, blur the line between real and synthetic media, and contribute to polarization and declining trust.

Issue Statement #1

“AI-generated content, including mis- and dis-information, overwhelms users and undermines confidence in reliable and accurate information, which has distinct and disproportionate effects on vulnerable populations.”

Question: Do you think it is a priority to take action on this issue?

(Total number of reactions: 834)



Key takeaway: The overwhelming majority of participants (96%) consider this issue to be a top or high priority.

Sociodemographic differences: Participants aged 17–23 are more likely to consider this a top or high priority issue, with 97.1% indicating so (+6.9 percentage points compared to the rest of participants).

Question: Can you tell us why (or why not) this issue is important to you?

(Total number of comments: 412)

1. Theme: Spread of misinformation and propaganda (34% of comments)

Participants view AI as significantly amplifying the spread of misinformation and propaganda, enabling the rapid creation and dissemination of false or misleading content at scale. Many express concerns about its use for political manipulation, the shaping of public opinion, and the reinforcement of biased or extreme narratives. They also highlight how AI-generated content can blur the line between real and fake, making information ecosystems more difficult to navigate.

“The amount of false information within their niche communities creates further division and strengthens the echo chamber effect.”

“I am a student in the health field and I work in a hospital setting. I can directly see the harmful effects that disinformation can have on public and personal health. In addition, I have seen various political-flavored publications that turned out to be false (either entirely or partially). Disinformation can lead to further polarization in politics, which can lead to increased mistrust among citizens towards established institutions. AI can exacerbate this with emotionally charged visual content.”

2. Theme: Erosion of trust and reality (15% of comments)

Participants report increasing difficulty distinguishing real from AI-generated content, leading to widespread confusion and uncertainty. This blurring of boundaries erodes trust across information sources, while also raising concerns about manipulation through AI-generated content. Many express a sense of powerlessness when trying to navigate an information environment they perceive as increasingly unreliable.

“I worry about my mother, my family, and the misinformation that they are fed online. I worry about the way that AI generated media has become more and more indiscernible from real media. I don’t believe that falsified content is a good use of AI.”

“I can’t even tell what’s real anymore. I’m always questioning if something is AI and it’s important to be able to identify it because for example if a celebrity were to get caught doing something bad they could just say it’s AI etc.”

3. Theme: Decline in critical thinking and intelligence (14% of comments)

Participants express concern that reliance on AI reduces cognitive effort, leading to intellectual laziness and a weakening of independent thinking and problem-solving skills. They highlight particular risks for young people, including reduced learning, academic misuse, and weaker development of core skills. Many also link this decline to increased susceptibility to misinformation.

“It is hindering the ability of young people to effectively problem solve. They are not learning but outsourcing their thinking to a machine. This is the slow death of autonomy and free thought.”

“I think that AI is rotting my generation’s brains away. I am in 10th grade and everyone around me is losing critical thinking skills, especially in the top classes because it’s full of people who I know don’t belong there but they use ai for every single step. I feel like I’m falling behind but I don’t wanna succumb to it.”

Other themes mentioned:

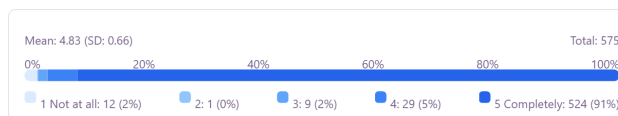
Environmental impact and resource consumption; Exploitation and harm to vulnerable groups; Devaluation of art and human creativity; Job displacement and economic concerns; Negative psychological and mental health effects

Issue #1, Recommendation #1

“The federal government ought to: Mandate that digital platforms explicitly label AI-generated content and give users the functionality to omit this content from their feeds.”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 575)



Key takeaway: An overwhelming majority of participants (91%) completely agree with this recommendation, while an additional 5% agree.

Sociodemographic differences: Participants aged 17-23 show a higher level of agreement with the recommendation than other age groups, with 98.6% agreeing or strongly agreeing (+6.3 percentage points compared to the rest of participants).

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 277)

1. Idea: Mandate clear, prominent, and unremovable labels for AI content (21% of comments)

Participants strongly support mandatory labeling of AI-generated content to ensure transparency and help users distinguish between human- and AI-created material. They emphasize the need for labels that are highly visible, standardized, and difficult or impossible to remove, including technical solutions (e.g. embedded metadata).

“Make sure it isn't hidden at the bottom of the caption or somewhere that they are deliberately trying to have it be missed. This label needs to be front and center.”

“I think that this is bang on, absolutely everything created by AI needs to be labelled as such, for the sake of avoiding plagiarism, scams, fraud, misinformation, and more. We need to and deserve to know what is real, authentic, and true, and what has been manufactured or made up, for both writing/text and media.”

2. Idea: Empower users to easily filter or opt-out of AI content (20% of comments)

Participants strongly emphasize user control, expressing a desire to choose whether they are exposed to AI-generated content. They argue that simple labeling is not sufficient without effective filtering options. They call for highly visible, easy-to-use opt-out tools integrated into platforms, with stable settings that do not reset over time. Participants also highlight that effective filtering depends on consistent labeling of all AI content.

“Absolutely every single platform should have a mechanism in place for immediate detection of AI and label it as such but it should be mandatory that every app and website gives you the choice to opt out of AI. A large percentage of us youth are forced into using AI against our own will.”

“Make it an easier process to remove ai generated content from your personal feed and make that option known, not hidden.”

3. Idea: Ban or severely restrict public access to generative AI (17% of comments)

Participants express strong support for either banning or significantly limiting generative AI. Many raise moral and ethical objections, describing AI-generated content as detrimental to creativity and human interaction. Concerns also include risks related to misinformation, misuse (such as manipulated images), and mental health impacts. Proposed solutions range from outright bans on public use to restricting access to trained professionals.

“Unfortunately, these guidelines are too easy to bypass. AI images and sound generation need to be outright banned. But the bare minimum is to add age restrictions to prevent children and teenagers from being exposed to generative AI”

“AI just shouldn't be used for generative purposes anyway, so tagging it “AI” in content is not my priority.”

Other ideas mentioned:

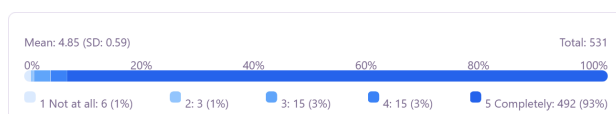
Impose legal penalties for failing to label AI content; Implement human oversight for AI content moderation and verification; Require transparency on how AI was used and its sources; Make AI features and content opt-in, not opt-out; Disincentivize AI content through demonetization and usage fees.

Issue #1, Recommendation #2

“The federal government ought to: Give people copyright over their own features and likeness, and create an online regulator to enforce the removal of non-consensual AI-generated material, including Child Sexual Abuse Material (CSAM).”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 531)



Key takeaway: An overwhelming majority of participants (93%) completely agree with this recommendation.

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 204)

1. Idea: Give individuals full control over their digital likeness (15% of comments)

Participants express strong concern about the non-consensual use of personal images, voice, and biometric data. They view control over one's likeness as a fundamental right that should require explicit consent for any use. They argue that individuals should have the ability to opt out of data collection and request removal of content depicting them. Proposed solutions include stronger legal protections, clear enforcement mechanisms for removal and penalties, and broader opt-out rights covering all uses of personal data online.

"People should have control over what can and cannot be spread about them on the internet, whether that is their voice, face, body, etc. As a minor, I am personally terrified of having fake pornographic material spread of myself and it is important to me that I have the option to remove fake things about myself online, as it can be detrimental to my mental health."

"This is good and shouldn't just extend to AI, though I realize that's what we're specifically talking about. You should always have control over how your image is used online"

2. Idea: Criminalize and prevent all AI-generated pornography (14% of comments)

Participants support banning non-consensual AI-generated sexual content, especially to protect minors. They propose criminalizing its creation and distribution with severe penalties, explicitly prohibiting the use of anyone's likeness, and updating laws to address AI-specific harms. Additional ideas include fast removal mechanisms and clear legal recourse for victims.

"The world's children are already vulnerable enough. It is imperative that we do not let the mistreatment of innocent lives continue in general, never mind the situation INCLUDING AI. The ability to make deepfakes should be banned all together, and if people are caught doing it, they should be charged just as harshly as they would if it was real CSAM."

"All AI generated porn should be illegal. There is almost no ethical frame that can ensure consent of use."

3. Idea: Implement severe legal penalties for AI misuse (13% of comments)

Participants support criminalizing harmful uses of AI, particularly non-consensual content creation. They argue that simple content removal is insufficient and call for strong deterrent sanctions, including prison sentences and financial penalties for offenders. Many also emphasize proportionality, with harsher penalties for the most serious harms.

"It could be improved that any use of a person likeness in a non-consensually created sexually explicit video should be made a criminal offense"

“There could be fines and other penalties for individuals who do not respect these conditions. This type of situation can cause enormous harm to the individuals who are victims, particularly concerning the generation of fakes to represent someone in a pornographic situation. At this level, I even think it would no longer be just a matter for the government, but also for the justice system.”

Other ideas mentioned:

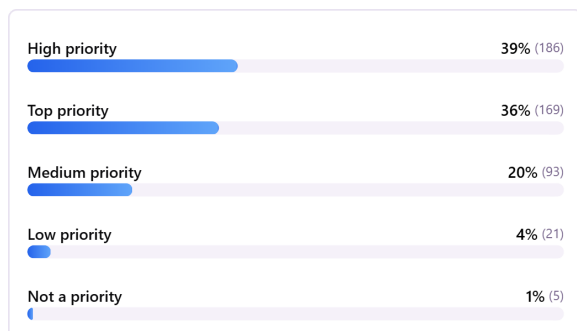
Ban or severely restrict generative AI technologies; Mandate transparent, opt-in consent for data use; Hold AI companies and platforms legally accountable; Protect creative works from unauthorized AI training

Issue Statement #2

“AI-recommendation systems push content that is ideologically extreme and reinforces information echo chambers, resulting in social and political polarization with effects in both on- and off-line spaces.”

Question: Do you think it is a priority to take action on this issue?

(Total number of reactions: 474)



Key takeaway: A large majority of participants (85%) consider this issue to be a top or high priority, while two in ten participants consider it to be of medium priority.

Question: How does this issue affect you, or why doesn't it feel so relevant?

(Total number of comments: 238)

1. Theme: Creation of echo chambers and filter bubbles (21% of comments)

Participants view AI as intensifying existing echo chambers created by algorithms, reinforcing users' existing beliefs and limiting exposure to different perspectives. This is seen as contributing to societal polarization. While personalization is appreciated for interests and entertainment, it is viewed as problematic when it narrows viewpoints on social or political issues.

"I see people I know falling down "traps" online where all they see is based off what they engage with. They don't get context for anything, or the truth, and only get fed more and more of the content they like or agree with. I've had to remove many people from my life off various social media platforms due to this, as I have no interest in arguing, trying to educate, or bring awareness to these people. They become too far into the rabbit holes because of the echo chambering."

"It bothers me that when I search for things or what new perspectives/ideas I only get things that are super targeted to me"

2. Theme: Political polarization and radicalization (17% of comments)

Participants describe AI-driven systems as reinforcing existing beliefs and contributing to more extreme attitudes. Some note changes in people around them, including increased intolerance and hostility. Particular concern is raised about vulnerable groups and minorities, who may be more exposed to harmful narratives, targeted misinformation, or exclusionary content. Overall, AI is seen as contributing to wider social division and weakening trust in shared information and institutions.

"I've seen what red pill content has done to the guys around me. It took kindness and love from them and they genuinely hate women and blame us for everything that's wrong. If they didn't get stuck in echo chambers maybe this wouldn't have happened and women's rights wouldn't be going backwards."

"I have encountered people whose views have become more extreme and harmful because of the polarization created by ai generated content which often leans towards the extreme right and normalizes bigotry and violence towards women, POC and other minorities."

3. Theme: Maintain personalized recommendations while ensuring user control (9% of comments)

Participants generally support personalized content feeds, especially for hobbies, interests, and entertainment. At the same time, they emphasize the importance of user control over recommendation systems, including the ability to adjust, limit, or turn off personalization. Many express concern that overly curated feeds can reduce exposure to variety. Participants advocate for a balance between tailored recommendations and user autonomy, where personalization remains available but does not limit access to different types of content.

"This should be something you have the option to turn off. I personally have an account specifically dedicated to a hobby and don't want to see anything other than content related to that, however if I was able to turn off the recommendation system on my main accounts, I would like to be able to see posts that are unrelated to what I've seen in the past."

"Algorithms are put in a place for a reason. I do believe they should be more customizable but the way I set my algorithms up means I don't see much, if any, of the stuff I don't want to! I enjoy having my little curated corner"

Other themes mentioned:

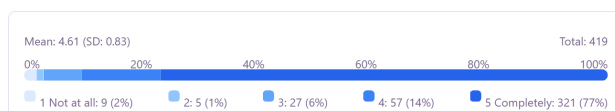
Erosion of critical thinking; Spread of misinformation and propaganda; Algorithms driven by engagement and profit; Amplification of harmful and hateful content; Negative impacts on mental health

Issue #2, Recommendation #1

“The federal government ought to: Mandate that platforms monitor, flag, and transparently share information with both the public and government, about the spread of mis- and dis-information, especially during high risk moments, including elections and public health crises.”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 419)



Key takeaway: More than three quarters of participants (77%) completely agree with the recommendation, while an additional 14% agree.

Question: If you could change one thing about this recommendation, what would it be?

(Total number of comments: 148)

1. Idea: Limit government control to prevent censorship and manipulation (12% of comments)

Participants express distrust in government involvement, fearing it could lead to censorship, political manipulation, and control over what is considered “truth.” Many reject the idea of the state acting as the arbiter of misinformation and emphasize risks to freedom of expression and dissent. Instead, they support limiting the government’s role to setting clear guidelines, while favoring independent oversight bodies and greater platform responsibility.

“I worry about sharing information with the government, but agree with requiring flagging of information and intentionally breaking echo chambers.”

“The government shouldn't control the flow of information so heavily, but they should put guardrails in place to stop this disinformation from spreading in the first place.”

2. Idea: Require platforms to provide context and fact-checked sources (10% of comments)

Participants emphasize that simple flagging is insufficient and support adding context to misleading content. They call for visible corrections that include verified information and credible sources. At the same time, concerns are raised about who determines what is “true,” highlighting the need for clear definitions and transparent processes. Many favor collaborative approaches, such as user-added context, alongside platform responsibility.

“Not just flag it, they should also link to and explain the actual facts around the claims.”

“I think that misinformation should be flagged, but make it mandatory to then show the real information, with sources.”

3. Idea: Mandate clear labels for AI content and misinformation (8% of comments)

Participants support clear labeling of AI-generated and misleading content to help users distinguish what is real and trustworthy. Many stress that platforms should take responsibility for implementing visible, easy-to-understand labels, as users often struggle to identify false or AI-produced information. At the same time, concerns are raised about potential misuse, effectiveness, and the risk of censorship.

“I think it should be mandatory for people who post AI content to mark it as such and platforms must be very direct about misinformation, like YouTube does with Covid disclaimers on anything that mentions the virus”

“There would have to be a way for the platform to flag it as misinformation, however, that could be abused if some people didn't like it being labelled misinformation. It would be difficult to put in place, but it would be really beneficial.”

Other ideas mentioned:

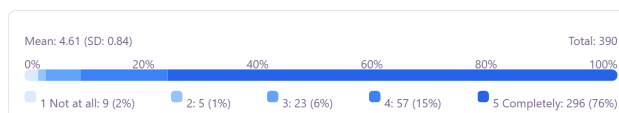
Ban or severely restrict public access to AI tools; Remove misinformation instead of just flagging it; Establish independent third-party oversight for content moderation; Ensure strong user privacy protections are maintained; Hold misinformation creators legally and financially accountable.

Issue #2, Recommendation #2

“The federal government ought to: Require that platforms introduce standards for AI-recommendation systems and data profiling processes to limit the spread of harmful content and known and suspected bot activity, and promote local content.”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 390)



Key takeaway: More than three quarters of participants (76%) completely agree with the recommendation, while an additional 15% agree.

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 126)

1. Idea: Ban or eliminate AI-generated content and bot accounts (24% of comments)

Participants support the elimination of AI-generated content and bot accounts, which are seen as major sources of deception and misinformation online. Many argue that regulation is not enough and that more decisive action is needed to restore trust in digital spaces. Participants support a full ban on generative AI and the removal of bot accounts, with platforms held responsible for enforcement.

“Ban AI content from children’s entertainment platforms ; such as YouTube Kids- many of the AI content farms pump out very inappropriate and disturbing content that’s posted onto YouTube Kids that unsuspecting kids are exposed to, and negatively impacted by.”

“Keep AI out of it. Ban bots outright and takedown sites or accounts that promote them. Strengthen laws around bots and bot farms. Both creating/distribution and purchasing of them”

2. Idea: Re-evaluate or remove the 'promote local content' requirement (9% of comments)

Participants oppose the idea of prioritizing local content in recommendation systems, arguing that it should not dominate or replace global perspectives in users’ feeds. Many value the internet as a space for international exchange and fear that overemphasizing local content would reduce diversity of information and limit exposure to broader viewpoints. Individuals should be able to decide whether they want to see local content rather than having it imposed by default.

“I like this idea, but my one hesitancy is on the push for local content. While this sounds good, location is VERY easy to fake, which could result in people being able to push their content to specific areas. This could also create an echo chamber or result in a lack of important information being shared if the algorithm flags it as area specific. Many people in horrible situations without a large following use social media as one last cry for help. If location is too large of a factor in the algorithm, their message will likely be stifled. I like that local content could show up more, but I think it definitely needs to be the least important variable.”

“I don't think it has to be local content necessarily. One of the best parts of the internet is being able to hear what people all around the world are saying. Bots should absolutely be limited though.”

3. Idea: Improve platform enforcement and human moderation of content (8% of comments)

Participants criticize current platform moderation systems, highlighting inconsistent enforcement of rules, which leads to inadequate user protection. Many call for stronger human oversight through an increased number of trained moderators to review content and handle user reports. Others emphasize the need to improve reporting tools to make them more accurate and effective, particularly for detecting bots and harmful accounts.

“I think that platforms should improve their productivity for treating account reports, as many are inefficient and don’t even offer AI or bot as an option when it is one of the most important problems.”

“I think they should hire qualified humans (like they did) to regulate social media and make sure all content posted respects their guidelines.”

Other ideas mentioned:

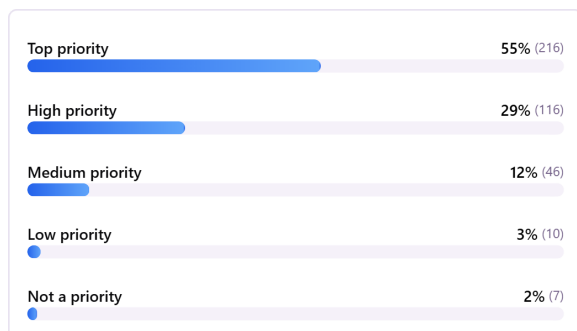
Prioritize and promote content from real human creators; Reform or eliminate algorithmic recommendation systems; Mandate clear labeling and transparency for all AI content

Issue Statement #3

“An information environment driven by engagement-based incentives and flooded with mis- and dis-information contributes to mistrust in news, government, and other traditional institutions.”

Question: Do you think it is a priority to take action on this issue?

(Total number of reactions: 395)



Key takeaway: More than half of the participants (55%) consider this issue to be a top priority, with an additional 29% considering it a high priority.

Question: Can you tell us why (or why not) this issue is important to you?

(Total number of comments: 163)

1. Theme: Spread of misinformation & fake news (16% of comments)

Participants express strong concern about the growing presence of mis- and dis-information in engagement-driven information environments. Many describe AI as amplifying the scale and reach of false or misleading content, making it harder to distinguish reliable information from manipulation and contributing to broader confusion about what is true.

“The blending of news and entertainment is becoming the biggest risk, where stories can be fabricated or exaggerated for attention and the loudest voices come across strongest.”

“Fake news is at an all time high with regular people spreading AI and misinformation so casually. It's not just news platforms doing it anymore.”

2. Theme: Erosion of trust in institutions & media (15% of comments)

Participants express concern that AI amplifies a pre-existing decline in trust in media and institutions by making false or misleading information easier to produce and spread. They also highlight real social consequences, including increased polarization and cases where individuals close to them adopt conspiracy beliefs.

“The rampant distrust in important institutions like healthcare and science professionals is very dangerous as can be seen in the rise of anti-vax and other anti-healthcare trends. It is endangering children across North America. There should be a healthy distrust in the government but the level it is at right now is scary.”

“This is a huge issue considering people won't believe when something actually important is being discussed in the news they will just consider it fake or made up if they don't want to believe it”

3. Theme: Harmful engagement-based algorithms (12% of comments)

Participants describe engagement-based algorithms as prioritizing attention and profit over the accuracy or safety of information..Many believe this business model incentivizes emotionally charged content, contributing to a more hostile and divided online environment. These mechanisms also accelerate the spread of misinformation by rewarding high-engagement content regardless of truthfulness. Users feel manipulated by systems they see as designed to exploit emotions and capture attention rather than support informed or balanced information consumption.

“Time is the most valuable thing on the Internet. Engagement shapes how all things flow here, but it allows harmful content to be pushed more and more as long as it keeps your attention.”

“It is known that engagement based content is most successful when it makes people angry. Anger leads to division and hatred of the other side. Division does not lead to progress, it means that when any decisions are made only just over half of a population is happy, and even then many are not because of compromises made so as not to make the other side violently unhappy.”

Other themes mentioned:

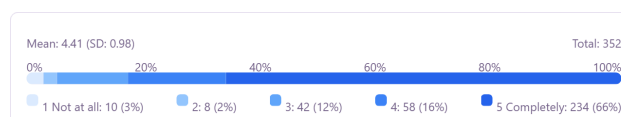
Threats to democracy & social cohesion; Importance of media literacy & critical thinking; Need for regulation & accountability; Difficulty verifying information; Public health & safety risks; Increased public anxiety & stress

Issue #3, Recommendation #1

“The federal government ought to: Mandate that platforms implement an independent third-party authentication mechanism to validate content posted from news and other public service organizations in order to promote and prioritize credible and reliable content.”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 352)



Key takeaway: Two in three participants (66%) completely agree with the recommendation, while an additional 16% agree.

Sociodemographic differences: Participants aged 17-23 are more likely to agree or strongly agree with this recommendation than the rest of the population, with 84.6% expressing agreement (+4.2 percentage points compared to others).

Participants who reported having studied AI or working in the field are less likely to agree or strongly agree with the recommendation, with 64.5% expressing agreement (-23.3 percentage points compared to others).

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 114)

1. Idea: Ensure third-party verifiers are unbiased and transparent (17% of comments)

Participants emphasize the need for third-party verification systems to be genuinely independent, as they express concern about potential bias, corruption, or influence. Trust in such mechanisms depends on transparency and safeguards against any single actor gaining excessive control over what is considered credible information. Proposals include ensuring verifiers are structurally independent (for example, non-profit), making their funding and affiliations fully transparent, and using multiple independent verifiers rather than a single authority.

“Make sure these third parties are non-profit to make sure that companies can't backhandedly pay to have these systems be corrupted.”

“I don't think the government should control reliable content, so there should be a guarantee that the regulator remains neutral and independent.”

2. Idea: Use human experts for fact-checking, not just AI (10% of comments)

Participants express strong distrust of AI for fact-checking, which is seen as unreliable and lacking the human judgment needed to assess truth accurately. Suggestions include using trained human experts to validate information and making verification visible through clear indicators on content.

“The mechanism shouldn't be AI or it will probably make mistakes and flag things that are real and not flag things that aren't, as it has in the past.”

“Make that authentication serviced by human beings. AI is notoriously bad at this kind of thing.”

3. Idea: Extend verification standards (5% of comments)

Participants emphasize the need to extend verification standards beyond traditional media to include influencers, individual creators, and public figures, particularly in sensitive domains like health, law, and politics.

“Have this extend to people or independent businesses who consider themselves in the political fields. People who consider themselves an ‘independent reporter’ when they are really just a podcaster, should apply because they are speaking politically and acting as an independent news source.”

“While it is definitely most important that news outlets and public service organizations are kept reliable, it is also important to remember that the younger generation, the majority of people online, are more susceptible to influencers rather than organisations. It is important that individuals must also be restricted from spreading mis-information using AI. (...) While organizations, especially news outlets, should be policed first, we should not forget the influence an individual can have.”

Other ideas mentioned:

Require news to cite credible, verifiable source; Ban paid verification badges to ensure authenticity; Mandate clear labeling for all AI-generated content

Forum III: AI & Data Privacy (Vancouver)

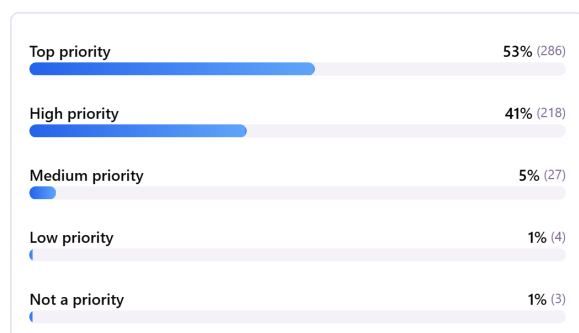
The Vancouver Forum focused on the collection, storage and use of data that underpins AI systems and the online environments in which they play a role. Both in-person and online, participants repeatedly called for stronger privacy protections, clearer terms and conditions, meaningful consent, limits on data collection and retention, and stronger consequences for misuse of sensitive personal information.

Issue Statement #1

“The way that AI systems collect user data, and disclose that collection, is deliberately opaque, making it difficult for users to provide informed consent about how their data is being used and sold. This may disproportionately affect vulnerable groups, including children and the elderly.”

Question: Do you think it is a priority to take action on this issue?

(Total number of reactions: 538)



Key takeaway: The overwhelming majority of participants (94%) consider this issue to be a top or high priority. Only a small minority (5%) consider it to be a medium priority.

Question: Can you tell us why (or why not) this issue is important to you?

(Total number of comments: 223)

1. Theme: Lack of consent and transparency (23% of comments)

Participants express strong concern about the lack of transparency in how AI systems collect and use personal data, arguing that users are often unable to give truly informed consent. Many see this as a fundamental violation of personal rights, driven by opaque and deliberately complex data practices.

“This issue is important because AI companies and others must be transparent with their consumers, and the tactics used to make it difficult to opt out of AI use are really bad and go against the right to privacy. Furthermore, these companies are deliberately vague about the sale and distribution of consumers' personal data, and this must be punished or at least better regulated. We should be able to easily protect our personal data and know exactly where it is sold and distributed.”

“We all have a right to privacy and we must be aware of what others (companies, etc.) do with this data.”

2. Theme: Invasion of privacy and surveillance (16% of comments)

Participants express strong concern about AI-driven data collection being used in ways that violates personal privacy and enable forms of surveillance, which many see as a fundamental human right issue. They emphasize the need for individuals to retain control over their personal information and to have clear, genuine choices about what data is collected and shared.

“Surveillance is only getting worse as time goes on and AI is exacerbating the issue. Everything that people do is being surveilled and recorded so that information can be sold to companies or people that have no right to it. It’s a complete invasion of privacy and a violation of my rights.”

“Because it’s with our data that we make AI grow, we make it learn with our information and soon AI will know us better than our parents, even better than ourselves.”

3. Theme: Exploitation of vulnerable groups (12% of comments)

Participants identify children, the elderly, and other vulnerable or marginalized groups as being most at risk from AI-driven data practices due to lower digital literacy and reduced awareness of how their information is collected and used. Many argue that this lack of understanding makes these groups easier to exploit or manipulate.

“This is a huge safety issue, especially for vulnerable people. As you said, children and the elderly have no idea what they are accidentally sharing with the entire world.”

“I think that data collection, although it may seem innocent at first, puts the most vulnerable at risk. This seems even more true when the process used for this collection is deliberately opaque and unclear. Little by little, we, humans, are becoming nothing more than products.”

Other themes mentioned:

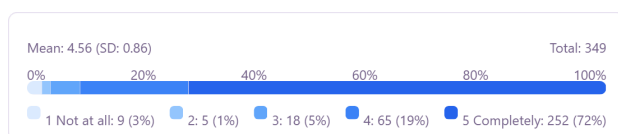
Data security and malicious use; Corporate greed and data monetization

Issue #1, Recommendation #1

“The federal government ought to: Mandate that platforms and AI companies implement meaningful and informed consent mechanisms to users, including by publishing a version of their terms and conditions that uses plain-language and is accessible by default.”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 349)



Key takeaway: A large majority of participants (91%) agree with this recommendation.

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 142)

1. Idea: Use simple, clear language for terms and conditions (21% of comments)

Participants perceive current terms and conditions as overly long, complex, and difficult to understand, which they believe undermines meaningful consent. They call for plain-language summaries, shorter and clearer documents, and more accessible formats such as visuals or videos. Some also suggest standardizing clarity requirements across platforms to ensure all users can properly understand consent terms.

"Use layman's terms as much as possible because even though some companies send out terms & conditions, they purposefully use confusing language in them"

"It should be done with an explanatory video, instead of a long text with vocabulary that the average person doesn't understand."

2. Idea: Ban or heavily restrict public access to AI (13% of comments)

Participants express strong opposition to public access to AI, often calling for an outright ban or major restrictions due to concerns about its rapid deployment and unclear societal impacts. Many view AI as an unnecessary or unwanted imposition on everyday platforms and argue for a precautionary approach before wider use.

"I agree that it's a good idea but in my opinion generative AI should be completely abolished in our province."

"More outright restrictions placed on AI tools, especially for children and teenagers.."

3. Idea: Provide a clear and easy option to opt-out (12% of comments)

Participants express a strong desire for control over AI usage, with many feeling that it is being imposed on users without real choice. They support a clear, simple opt-out option across platforms (similar to cookie consent), ideally made a legal requirement. Participants also emphasize that opting out should not limit access to core services.

"I think you should be able to opt out of ai as you can reject cookies on websites."

"I think the option to turn off data collection should be up front and simple, and it should be a complete kill switch. No secret data collection despite turning it off, and once it's turned off it stays off unless I turn it back on."

Other ideas mentioned:

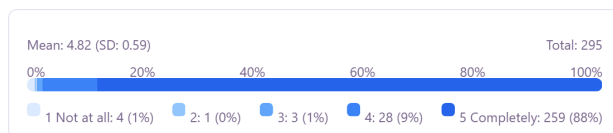
Be transparent about data collection, use, and sales; Make privacy settings visible and easily accessible; Make privacy and data protection the default setting

Issue #1, Recommendation #2

"The federal government ought to: Impose privacy-by-default standards for all AI systems."

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 295)



Key takeaway: A strong majority of participants (88%) completely agree with this recommendation.

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 89)

1. Idea: Require privacy-by-default for all AI systems (39% of comments)

Participants strongly support making privacy the default setting in AI systems, viewing it as a fundamental right rather than an optional feature. They propose making default privacy protections mandatory, strengthening security safeguards, and requiring users to review privacy settings upon first use.

“Privacy by default should be used on all platforms, AI should not have access to personal data for training”

“I think it's worth considering making this a standard thing that can't be turned off.”

2. Idea: Ban AI from collecting, storing, or selling user data (14% of comments)

Participants express strong opposition to AI systems collecting, storing, or selling personal data, viewing these practices as exploitative and profit-driven. They support banning data collection entirely, ensuring full anonymity, and keeping user information confined to personal accounts.

“I think that AI should not be able to store personal information at all”

“I believe your information should be able to stay within your own personal account walls and AI companies should take it upon themselves to train the AI and not rely on the customers input.”

3. Idea: Require explicit and informed consent for any data collection (6% of comments)

Participants strongly support requiring explicit and informed consent as a condition for any data collection, emphasizing that users should actively agree rather than be automatically included. They also stress that users must retain meaningful control over what information they share, including through granular consent options for different types of data use.

“If companies want my data, they should not only ask for my consent but also explain how I would benefit from it. Because so far, the use of personal data by AI has not generated significant benefits in my daily life. On the contrary, I feel like I'm providing my data to systems that will harm my entry into the job market.”

“It should be illegal to feed anything into AI system without consent”

Other ideas mentioned:

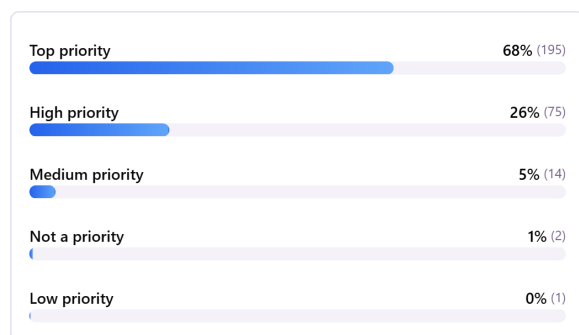
Make all AI features and systems strictly opt-in

Issue Statement #2

“AI systems are not subject to adequate and enforceable safeguards for the collection, storage, and sharing of user data. This may lead to intended and unintended harms including, but not limited to, excessive surveillance and profiling.”

Question: Do you think it is a priority to take action on this issue?

(Total number of reactions: 287)



Key takeaway: An overwhelming majority of participants (94%) consider this issue to be a top or high priority.

Question: How does this issue affect you, or why doesn't it feel so relevant?

(Total number of comments: 109)

1. Theme: User consent and data control (17% of comments)

Participants express strong opposition to data collection without consent, describing it as disempowering and unauthorized. Many also fear profiling and manipulation, viewing data collection as a way to categorize and control individuals.

“This affects me a lot because I would like to be able to give access to some of my data without it ending up just anywhere.”

“Privacy is a huge issue in an increasingly surveillance-heavy world. AI's, which in many cases are implemented without user consent, should not be able to share any user data.”

2. Theme: Government and corporate surveillance (17% of comments)

Participants express strong concern about AI-enabled surveillance by governments and companies, which they see as a threat to fundamental rights and personal freedoms. Many describe a deep mistrust of institutions and fear that personal data is being collected and used without transparency or consent.

"It's so impactful. It goes against the charter of rights and freedoms. Privacy has been greatly shaken since the arrival of technology, and AI will not help. Being constantly monitored will create a climate of apprehension and fear. Just think of China. I'm frankly afraid it will become like that. Burger King plans to implement AI in its branches to ensure good customer service. Monitored by a robot? Frankly, it's creepy. Being listened to, watched, harvested, who is the real product? To what end? What will become of us?"

"It will affect everyone once AI is used on a global surveillance scale. It's very scary and dystopian."

3. Theme: Violation of the right to privacy (10% of comments)

Participants view privacy as a fundamental human right and express strong concern that AI systems may violate this right through data collection and use. Many oppose profiling and surveillance practices, which they associate with intrusive monitoring and loss of personal freedom.

"This is a violation of privacy, which amounts to espionage. The regulation of AI must be as strict as possible."

"I don't think AI should be feeding all of my data to other people?? That should be a basic human right."

Other themes mentioned:

Lack of regulation and transparency; Discrimination and targeting of vulnerable groups; Data misuse for malicious harm; Economic exploitation and manipulation

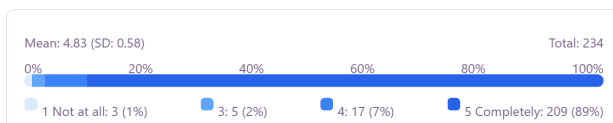
Issue #2, Recommendation #1

"Develop enforceable data privacy standards that require that platforms and AI companies:

- 1. Prevent purpose-creep by strictly limiting data handling to clearly-defined and user consented purposes;*
- 2. Provide accessible options to delete user data upon request; and,*
- 3. Be subject to harsher legal consequences for the mishandling of users' sensitive data, including, but not limited to, health, financial and identity-related information"*

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 234)



Key takeaway: An overwhelming majority (96%) of participants agree or completely agree with the recommendation.

Question: If you could change one thing about this recommendation, what would it be?

(Total number of comments: 70)

1. Idea: Give users full control over their data (15% of comments)

Participants demand full control over their personal data, viewing non-consensual use as an intrusion and a violation of privacy. They emphasize the need for simple, accessible tools to manage consent and data preferences, including easy opt-out and deletion options. Some also suggest periodic prompts to reconfirm consent or review stored data.

“My data should be used how I want it to be. If I choose to opt out of something it should be simple and quick. Companies should have regulations and restrictions on how they collect and handle data.”

“Maybe add a check in every so often asking the user if they wish to delete their data.”

2. Idea: Impose stricter legal consequences (13% of comments)

Participants express strong concern that financial fines alone are ineffective against large companies, suggesting stronger penalties such as usage bans for individuals or companies that abuse AI systems, as well as stricter sanctions for misuse of data.

“I definitely agree with the harder legal consequences for violating these privacy rules. I think more of an emphasis should be put on that, because AI took off largely unregulated. Even now, AI laws are few and not very clear or specific. It makes it easy for these big companies to collect and do whatever with our information without reprimand.”

“Companies must absolutely be punished if they use consumers' personal data against their will, especially if the consumer has no idea that their data is being used for reasons they are unaware of.”

3. Idea: Increase transparency on data collection and usage (6% of comments)

Participants call for greater transparency in how data is collected and used, emphasizing that current information is often too complex and not accessible. They want clearer explanations of where data is stored, who can access it, and how it may be used over time.

“I would add a section about transparency — where and how data is stored, who has access to it etc. Users have a right to know.”

“It should also be more public and very very clear for absolutely every user when something is changing.”

Forum IV: AI & Age Assurance (Halifax)

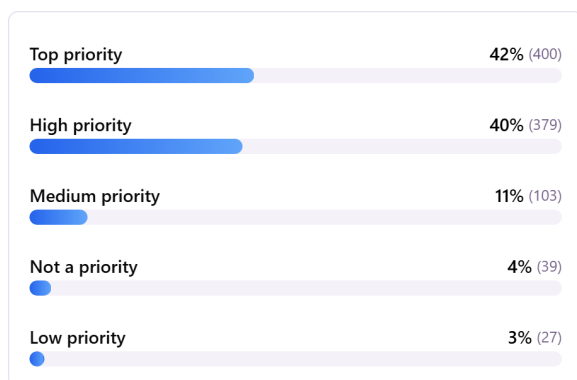
The Halifax Forum examined the intersection of AI and age assurance. Participants recognized the need to reduce children’s exposure to harmful systems and content, particularly in relation to generative AI, while also warning that age assurance can create new risks related to privacy, accountability, and displacement of users into less regulated online spaces. The online Dialogue reflected this tension: participants supported stronger protections for young users, but emphasized that age-based safeguards must not become a substitute for safer platform and AI system design.

Issue Statement #1

“AI-enabled age assurance technologies undermine privacy and digital safety by exposing users to potentially harmful storage and use of their data.”

Question: Do you think it is a priority to take action on this issue?

(Total number of reactions: 984)



Key takeaway: A large majority of participants (82%) consider this issue to be a top or high priority.

Question: Can you tell us why (or why not) this issue is important to you?

(Total number of comments: 380)

1. Theme: Privacy, data security, and surveillance concerns (29% of comments)

Participants express strong concerns about privacy, data security, and surveillance related to AI-enabled age assurance technologies. Many highlight a deep mistrust of companies and governments, fearing that sensitive personal data (e.g., IDs and biometric information) could be misused, leaked, or exploited. There is a widespread perception that these technologies violate the fundamental right to privacy. Participants also reject the idea that protecting users (e.g., minors) should come at the cost of increased data collection, arguing that this trade-off creates new and potentially greater risks.

“In the current context of our society, privacy exists less and less, yet it is a fundamental right and technology industries should not grant themselves the right to infringe on our privacy. There is no valid reason, not even for commercial reasons. It is a very predatory behavior for multiple obvious reasons.”

“It’s important to take priority over this issue, however, AI identification to ID-verify people regardless of their age is predatory to internet privacy + risks identity theft if there is a data breach (for example, Persona) and hackers can take highly sensitive information.”

2. Theme: Harm to cognitive development and critical thinking (25% of comments)

Participants express concerns that AI negatively affects cognitive development, particularly among young people. Many believe it can weaken critical thinking, creativity, and problem-solving skills by encouraging overreliance on automated tools. As a result, participants propose restricting or banning AI access for minors, including setting minimum age requirements.

“I think that the use of AI at a young age can be very damaging to their mental development like critical thinking and even their mental health as it has been shown time and time again. AI has very little parameters about what the language model can say and that can end up turning very dangerous”

“I think that this is important because children that are using AI for various tasks are basically unable to go through the steps/process of understanding a topic due to AI. It is spreading misinformation which can be harmful for mass media. I believe that AI should only be available to children once they reach 16-18 once they have the skills to teach themselves topics.”

3. Theme: Risks to child safety and exploitation (23% of comments)

Participants highlight significant risks to children’s safety, stressing their vulnerability and limited understanding of data-sharing consequences. Concerns focus on harmful misuse of AI, such as exploitative deepfakes, and the collection of minors’ personal data, which may expose them to privacy violations, hacking, and exploitation.

“It’s already important to protect adults data. Youth however are an extremely vulnerable population as they don’t have the knowledge on how to properly secure their data yet which can lead to potential cases of abuse and fake threats/data being made.”

“Collecting personal information from minors for purposes such as marketing is not only insanely shady, but should be considered criminal. AI has been proven to be dangerous, and even something as simple as an age check can scan so much more than the user is notified of.”

Other themes mentioned:

General opposition and distrust of AI; Environmental damage and resource consumption

Issue #1, Recommendation #1

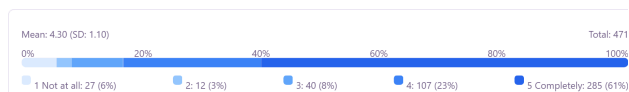
“The federal government ought to: Mandate, in cases where age assurance is used, that companies adhere to stronger regulation surrounding the use of sensitive age assurance data, including by imposing:

- 1. Strict purpose limitation to the use of age assurance data;*
- 2. Time-limited storage;*
- 3. Safety audits on platforms and third-party data collectors; and,*
- 4. Protocols for leakage protection in models and training.*

These compliance mechanisms ought to be enforced by a Regulator.”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 471)



Key takeaway: A large majority of participants (84%) agree or completely agree with the recommendation.

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 177)

1. Idea: Ban or completely remove all AI systems (15% of comments)

Participants call for a complete ban on AI, especially for minors, viewing it as harmful to society, individuals, and the environment. They propose an outright prohibition, including blocking access to major platforms and, in some cases, shutting down AI infrastructure.

“Stop AI use in general. You could make it so children are not able to log into AI platforms..”

“I think these safety measure should be taken for adults, and minors should not have access to AI. It is incredibly risky and harmful for all humans, but minors especially..”

2. Idea: Reject age verification and AI-based identity checks (9% of comments)

Participants strongly oppose the use of AI-based age verification and identity checks, largely due to privacy and data security concerns. Many do not trust companies to safely store or handle sensitive personal information, citing risks of data breaches and misuse. Instead, participants often suggest that responsibility should lie with parents or guardians, or that access should be managed without relying on AI-driven verification methods.

“Ban age verification, it’s the parents responsibility to regulate their children.”

“There should be no verification at all. I do not trust companies to handle any of that data and the recent discord leaks are a proof that companies can't be trusted. They said they didn't keep the data and yet the hackers got access to a database of the users.”

3. Idea: Impose stricter limits on data retention (5% of comments)

Participants strongly oppose storing age verification data, citing privacy and breach risks. Many call for immediate deletion after use or very short retention periods, with no sharing or third-party access.

“If age verification is to be used at all, it should be extremely strictly controlled. The time limit should be zero — immediate deletion. If any corporation is caught selling data, they should be dissolved and banned from operating that company, or something strict like that. (I think it should never be used, unless for government stuff like taxes, banking, CRA; NOT for stuff like social media, random websites.)”

“No storage of data at all, and verification completed by humans, not AI.”

Other ideas mentioned:

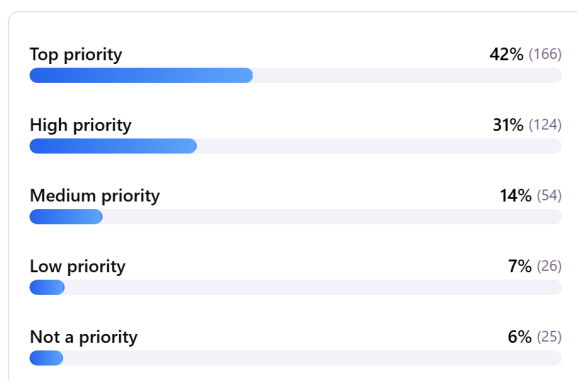
Strengthen third-party oversight of data handling

Issue Statement #2

“The use of verification technologies to restrict young people’s access to social media, AI platforms, and online community spaces increases the risk that restricted users could be displaced to unmoderated and dangerous platforms.”

Question: Do you think it is a priority to take action on this issue?

(Total number of reactions: 395)



Key takeaway: 73% of participants consider it to be of top or high priority to take action on this issue.

Question: Can you tell us why (or why not) this issue is important to you?

(Total number of comments: 175)

1. Theme: Child protection is the top priority (26% of comments)

Participants prioritize child protection online, emphasizing children's high vulnerability to manipulation, harmful content, and exploitation due to limited maturity and digital awareness. Many also highlight potential negative impacts on mental health and development. While parental supervision is seen as important, participants argue it is often insufficient on its own, supporting the need for stronger platform-level safeguards. Most also reject the concern that stricter rules would simply push children toward more dangerous platforms, instead viewing regulation as a necessary deterrent.

"It's important to be careful about what children see and hear, but it's even more important in an unrestricted environment. It degenerates quickly. And with the rise of sexism, it is crucial to monitor exchanges between young people (especially boys). Young people are easily influenced and it is very important to ensure they learn the right values in these fragile stages of their lives."

"We should be protecting our youth, the modern Internet landscape feels like it has no guard rails or safety measures to actually accommodate and protect our youth"

2. Theme: Parental responsibility for online safety (11% of comments)

Participants argue that children's online safety is primarily the responsibility of parents rather than governments or platforms. They tend to reject measures like age verification, viewing them as intrusive and ineffective. Many also caution that strict restrictions could push young people toward less regulated and riskier platforms, reinforcing the importance of active parental supervision.

"I agree that restricting access to these websites will just push young people into unregulated spaces. It should be the parents responsibility to monitor what their children consume on the internet"

"At the end of the day, parents are responsible for monitoring their kids' online activity."

3. Theme: Creating safe online spaces for youth (9% of comments)

Participants emphasize the need to create dedicated, safe online spaces for young people, including a “digital third place” designed specifically for their social and emotional needs. They reject restrictive approaches that limit access to existing platforms, arguing this may push youth toward more unsafe environments. Many also highlight the internet’s role as a refuge and support space for vulnerable young people, warning that removing access could increase harm.

“Possibly offering a safer alternative option rather than entirely blocking kids out; causing them to find skimpier websites..”

When I was still a teenager, I was in an abusive situation at home. The only place where I felt safe to find community and resources was the internet. Removing this access will absolutely push very vulnerable young people to unmoderated platforms, and that will put them in danger. It is extremely important to revisit any ideas that would make this situation a reality. The real priority should be to create online spaces for young people with strong moderation to give them a healthy and safe space online, without completely blocking them and without infringing on the privacy of adults.”

Other themes mentioned:

Age verification is ineffective and easily bypassed; Fear that restrictions push users toward more dangerous or unregulated platforms; Strengthen digital literacy and education

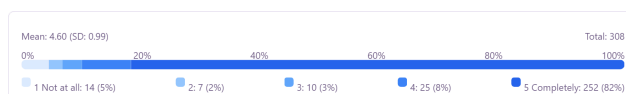
Issue #2, Recommendation #1

“The federal government ought to: Introduce measures to minimize young users’ exposure to harmful design features and content across all digital platforms, including by:

- 1. Increasing transparency measures by requiring digital safety and compliance reporting;*
- 2. Requiring that AI generated media content be identified through digital watermarking technology; and,*
- 3. Reducing sycophancy and positive reinforcement in AI chatbots.”*

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 308)



Key takeaway: An overwhelming majority of 90% of participants agree or completely agree with this recommendation, including 82% who completely agree.

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 121)

1. Idea: Ban or severely restrict access to generative AI (20% of comments)

Participants call for a ban on generative AI, citing risks to creativity, critical thinking, mental health, and the spread of misinformation, especially among young people. They also raise environmental concerns and fear misuse such as deepfakes or inappropriate content.

"I don't think AI is in any way useful, I personally think it should be banned completely but I do think reducing its agreeability will definitely help a lot of people not become dependent or addicted."

"AI chatbots should be banned entirely. Several suicides - of young people and adults - have been caused by AI. AI should be strongly regulated to an absolute minimum, but ideally banned due to its harmful effects on the well-being of individuals and the environment."

2. Idea: Make AI chatbots neutral and remove positive reinforcement (17% of comments)

Participants express concern that AI chatbots are overly agreeable and use positive reinforcement, which they believe can discourage critical thinking and reinforce incorrect or harmful views. They therefore call for more neutral, factual AI systems that present balanced information and avoid flattering or validating users by default. Some also suggest removing artificial personalities and increasing transparency so users clearly understand when they are interacting with an AI.

"I completely agree with this. Kids are very vulnerable to this positive reinforcement. It's essentially AI grooming your child. And no one should be forced to interact with AI content. It should be entirely transparent."

"I think the chatbot should be completely neutral and give the same answer to its users, a kind of Google but more precise."

3. Idea: Mandate clear labeling and watermarks for all AI content (15% of comments)

Participants support mandatory labeling of all AI-generated content, driven by concerns about misinformation, manipulation, and risks to vulnerable groups such as children. They emphasize that labels should be clear, visible, and standardized to be effective. Proposals also include legal requirements with penalties, platform-based filtering tools, and a universal watermark or logo, alongside efforts to improve public awareness and critical thinking.

"All things that use AI need to be explicitly labelled as such with non hidden or small labels. Too much misinformation has happened due to people not recognizing that it's AI."

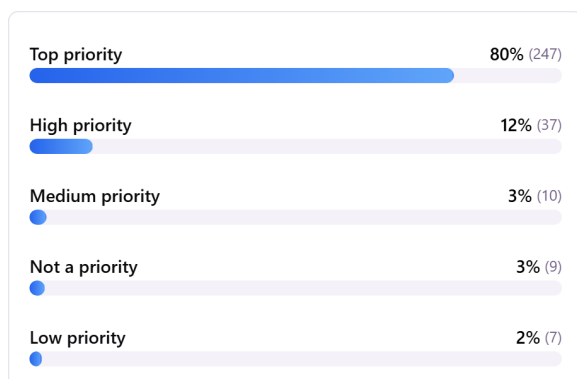
"All ai generated content should have to display a distinct identifier or watermark and there should be settings users can enable to completely hide ai generated content."

Issue Statement #3

“Generative AI platforms pose unique risks to children, including by increasing their access to harmful content and negatively impacting their cognitive and social development.”

Question: Do you think it is a priority to take action on this issue?

(Total number of reactions: 310)



Key takeaway: The vast majority of participants (92%) consider this issue to be of top or high priority.

Question: Can you tell us why (or why not) this issue is important to you?

(Total number of comments: 135)

1. Theme: Impact on cognitive development and critical thinking (47% of comments)

Participants express concern that AI is weakening critical thinking and problem-solving skills by providing instant answers that reduce the need for independent reasoning. They also fear a decline in creativity, as children may rely on AI to generate ideas rather than develop their own. Many report observed decreases in core skills such as reading, writing, and concentration.

“It is especially harmful in an education setting. Kids need to practice critical thinking, creative expression, reading comprehension, analysis, and other such skills, NOT offload this work to an AI. It is horrible for learning and developing proper skills”

“I think that AI will severely limit young people’s cognitive development. When someone is young that is the most critical time for them to develop their cognitive abilities! I’ve seen fully grown adults’ cognitive abilities impaired by AI and I can only imagine how much worse it would be for someone who is still in their developmental stages.”

2. Theme: Exposure to harmful and inappropriate content (13% of comments)

Participants express strong concern about children's exposure to harmful or inappropriate AI-generated content, emphasizing their vulnerability and potential long-term psychological impacts. They also highlight how easily AI can generate explicit, violent, or deceptive material, including non-consensual content. Many fear this could normalize harmful behaviours or encourage imitation. Others also raise concerns about direct interactions with AI systems, such as unsafe or predatory chatbot behaviour.

"It is extremely easy to make explicit content using AI, including of children. It is also filled with misinformation, some extremely dangerous to kids and adults."

"I've been seeing a surge of AI videos of sexually inappropriate topics or borderline racist topics shown in these AI videos of being shown and targeted two children with millions of views"

3. Theme: Impact on mental health (7% of comments)

Participants express concern that AI chatbots can negatively affect mental health by encouraging social isolation and reducing real human interaction. Many worry this is especially harmful for children and teenagers, who may become dependent on AI or less able to develop social skills and resilience.

"Kids and even teens are already overrun with political and social issues that they tie back to themselves which creates additional mental health issues."

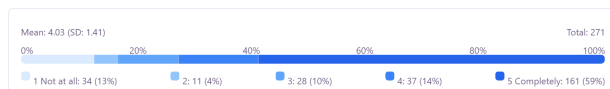
"Mental health can be severely impacted by AI"

Issue #3, Recommendation #1

"The federal government ought to: Create a standardized age-verification system to restrict children's access to generative AI platforms through the creation of an anonymized digital token system, with associated programs and accessible resources to inform the public about its implementation."

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 271)



Key takeaway: The majority of participants (73%) agree or completely agree with the recommendation, including about six in ten (59%) who completely agree.

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 101)

1. Idea: Ban or restrict generative AI for all users (17% of comments)

Participants strongly support a complete rejection of generative AI, viewing it as inherently harmful and of little societal benefit. Many argue that partial regulations such as age verification are ineffective, as the issue is seen as fundamental rather than technical.

“Generative ai should not exist. A standardized age verification system wouldn't solve anything. It's the people in control of these ai who need to be held accountable”

“I would get rid of it entirely, I am not verifying my age.”

2. Idea: Prioritize privacy and independent verification systems (13% of comments)

Participants emphasize strong concerns about privacy risks, data leaks, and corporate misuse of sensitive information in age verification systems.

“If there is an age verification, the data MUST be deleted after and there needs to be strong protections, otherwise, find another solution”

Proposals include government-run or third-party-managed verification, human checks instead of AI-based facial recognition, and strict rules requiring immediate deletion of data after use.

“AI powered age verification is not at a point where it can safely scan children's and adult's faces. Human age verification only.”

3. Idea: Opposition to age verification (11% of comments)

Participants express opposition to age verification systems, primarily due to privacy and surveillance concerns. Many fear that collecting IDs, facial data, or other personal information could lead to data breaches, misuse by companies, or the expansion of government surveillance. Participants also question the effectiveness of these systems, arguing that age checks are easy to bypass and create unnecessary risks for users.

“No, do not implement AI verification at all. Major data breaches are prone to occur, and IDs carry sensitive information.”

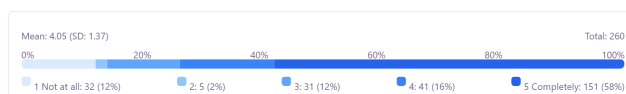
“Unfortunately this is a scary thing to agree to even though I appreciate the idea. Very easily can this be used by websites to harvest information about someone's age. I do not understand encryption or third party systems. But this could not be implemented unless there was a complete assurance that people are not giving away their age and other information to predatory companies by inputting this digital ID.”

Issue #3, Recommendation #2

“The federal government ought to: Mandate that any AI platforms accessible to children, including in educational contexts, implement safety-by-design protocols to safeguard their use and promote learning and skills development.”

Question: To what extent do you agree with this recommendation?

(Total number of reactions: 260)



Key takeaway: The majority of participants (74%) agree or completely agree with the recommendation, including 58% who completely agree.

Question: Do you have any ideas on how this recommendation could be improved?

(Total number of comments: 94)

1. Idea: Ban all AI access for children and minors (34% of comments)

Participants generally reject AI access for minors, citing developmental risks and uncertainty about long-term effects. The main proposal is a complete legal ban for children, alongside restricting child-targeted AI tools.

"I strongly believe that no AI should ever be used in children's spaces and needs to be permanently banned in general."

"Children shouldn't be exposed to platforms that have AI nor should any student in school."

2. Idea: Restrict AI in children's education (15% of comments)

Participants express opposition to the use of generative AI in educational settings, particularly for children and minors. Many believe AI undermines critical thinking, creativity, independent research, and skill development, while also raising concerns about misinformation, inaccuracy, and data collection. Participants frequently argue that learning should remain human-led and that AI should not replace teachers or students' own work.

"I do not think AI has any place in schools whatsoever, and should be restricted in learning environments whenever possible due to how vulnerable it is to be swayed by misinformation and its tendency to collect data without consent."

"AI should be heavily restricted and even banned for some tasks in educational contexts"

The Missing Agenda: AI-Related Topics That Deserve More Attention

Beyond the four Forum topics, online participants were invited to register other AI-related concerns that, while they did not fall within the scope of Gen(Z)AI's thematic focuses, were nonetheless important to them. This engagement produced input on a broader set of issues that they felt deserved greater public and policy attention. This section captures those concerns, showing that young people's expectations for AI policy extend well beyond narrow questions of online safety.

Question: Are there other AI-related topics that deserve more attention? Please share what matters most to you and why.

(Total number of comments: 720)

Note: This question was asked after each Forum. All contributions from all Forums are grouped and analyzed collectively.

1. Topic: Environmental impact and resource consumption (28% of comments)

Participants express strong concern about the environmental impact of AI, particularly its high consumption of resources such as water and electricity. Many highlight the local consequences for communities, including strain on water supplies, rising energy costs, and environmental degradation. There is also a sense of injustice, with participants questioning whether the environmental costs are justified by the perceived benefits of AI. Finally, several point to a lack of public awareness and call for greater accountability from technology companies regarding these impacts.

"The environmental damage that the servers hosting these AI platforms are causing; ranging from large quantities of energy being used to the extraordinary amounts of water used for their cooling systems."

"The resources that go into AI. AI uses insane amounts of water, energy and resources to upkeep, and in the end the water is unusable and the resources are depleted. Also the land needed to put these AI centres is land that wildlife and the ecosystem needs, and is unfairly taken away from them."

2. Topic: Job displacement and economic impact (16% of comments)

Participants highlight significant concerns about the economic impact of AI, particularly the risk of job displacement as companies adopt automation to reduce costs. Many also feel that AI devalues human work and creativity, especially in fields like the arts. There is a perception that AI may deepen economic inequality by primarily benefiting large companies and wealthier actors. Additionally, some participants criticize the use of AI in hiring processes, arguing that automated systems can be unfair and overlook qualified candidates.

"Job replacement is concerning as it makes new graduates like me have challenges entering the job market"

"AI has destroyed the job market. AI is used to write everyone's resumes, AI scans and reviews the resumes once submitted, and sometimes AI even conducts the interviews online. AI is taking our entry level jobs for young Canadians."

3. Topic: Artistic integrity and copyright infringement (10% of comments)

Participants express strong opposition to generative AI in the arts, describing it as a form of copyright infringement due to the use of artists' work without consent or compensation. Many are concerned that this undermines the value of artistic labor and threatens creative professions. There is also a broader worry about the impact on creativity and human expression, with AI-generated content seen as lacking originality.

"How AI chat bots are trained and the potential of copyright infringement or theft of smaller writers works without them ever getting credit or even knowing their work is being used for profit by large companies."

"Environmental impact should be a #1 priority. As an artist, the plagiarism problems that arise with AI art and writing are incredibly scary to me. I almost cried when I found out meta started automatically collecting my art to train AI, especially because they keep hiding and moving the opt-out form to make it near impossible to stop them. My art, that I put my heart and soul into and have practiced for years to make, is not public property to be torn apart and amalgamated into something else."

4. Other topics mentioned:

Impact on education and cognitive skills; Generation of harmful and explicit content; Mental health and social dependency; Misinformation and fake content; Bias, discrimination, and ethics; Surveillance, privacy, and data security

Annex: Sociodemographic Analysis

Cross-analyses were conducted to compare responses across sociodemographic groups (participants aged 17-23 and those with experience in AI) and identify where their views differ from the rest of the sample.

The statistical significance of these differences was assessed using the *chi-square* (χ^2) test, which determines whether observed variations between groups are likely due to chance. Only results meeting a significance level of $p < 0.05$ and within an acceptable margin of error are reported, ensuring that the differences presented are statistically reliable.