



Public evaluations of AI-mediated healthcare when considering people with diverse access and support needs: a mixed-methods survey

Joni-Roy Piispanen^{a,*}, Janne Kauttonen^b, Ari Alamäki^b, Rebekah Rousi^a

^a School of Marketing and Communication, University of Vaasa, Wolffintie 32, 65200 Vaasa, Finland

^b Haaga-Helia University of Applied Sciences, Ratapihantie 13, 00520 Helsinki, Finland

ARTICLE INFO

Keywords:

Artificial Intelligence
Healthcare AI
Accessibility
Disability
Health Equity
Patient-Centered Care
Algorithmic Fairness

ABSTRACT

Background: Artificial intelligence (AI) is increasingly being introduced into healthcare, including patient communication, monitoring, triage, decision support, and care-related services. Accessibility is central to these developments, particularly when AI-mediated healthcare is considered in relation to people with diverse access and support needs.

Objective: This study examined how accessibility-related evaluations of AI in healthcare are patterned when respondents are prompted to consider people with diverse access and support needs, and how respondents describe risks and conditions of acceptable AI-mediated care in this context.

Methods: We conducted a mixed-methods questionnaire study with a panel sample (N = 1,151). Quantitative analyses combined Bayesian regression, confirmatory factor analysis, and correlation analysis to examine accessibility-related evaluations of healthcare AI. Qualitative analysis examined open-ended concern elaborations from a subgroup of respondents (n = 117) using codebook thematic analysis with collaborative coding.

Results: Support-oriented evaluations clustered together more strongly than they aligned with concern. In the selected regression models, technological acceptance and age were the clearest correlates: higher technological acceptance was associated with more supportive evaluations, while older age was associated with greater concern and lower endorsement of AI's supportive potential in some models. Four qualitative themes were identified: Human Recourse and Relational Care; Communicative Accessibility and Misunderstanding; Safety, Reliability, and Override Capacity; and Privacy, Misuse, and Unfair Classification. Respondents articulated concern through questions of relational care, communicative fit, accountability, fairness, safety, and access to human support when AI systems fail or are misunderstood.

Conclusions: The findings suggest that healthcare AI implementation should address intelligibility, communicative accessibility, bounded system roles, accountability, and human recourse as central conditions for acceptable AI-mediated care.

1. Introduction

Artificial intelligence (AI) is being introduced into healthcare in areas such as patient communication, triage, monitoring, diagnosis, and decision support [1,2]. Recent scholarship has established accessibility and inclusivity as important issues in AI-enabled healthcare and health-related social care [3,4]. Research focused more directly on disability and access shows that AI can expand participation and can also constrain it, depending on how systems are designed, trained, and deployed [5–8]. Related work suggests that AI literacy, digital confidence, and access to

devices, connectivity, and support shape who is able to benefit from AI systems and under what conditions [9–12]. An overview of what was previously known and what this study adds can be seen in Table 1.

Patients and members of the public have been found to often express broadly positive views about clinical AI, while also reporting reservations and a preference for human supervision [13]. Public surveys in the United States have examined healthcare AI attitudes across issues such as trust, perceived benefit, privacy, and acceptability, indicating that public support is often conditional on how AI is used and governed [14]. Attitudes toward diagnostic AI vary across patient groups and are

* Corresponding author.

E-mail addresses: joni.piispanen@uwasa.fi (J.-R. Piispanen), janne.kauttonen@haaga-helia.fi (J. Kauttonen), ari.alamaki@haaga-helia.fi (A. Alamäki), rebekah.rousii@uwasa.fi (R. Rousi).

<https://doi.org/10.1016/j.ijmedinf.2026.106625>

Received 27 April 2026; Received in revised form 22 June 2026; Accepted 16 July 2026

Available online 17 July 2026

1386-5056/© 2026 The Author(s). Published by Elsevier B.V. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

Table 1
Summary Table.

What was already known	What this study adds
AI is increasingly discussed and implemented in healthcare communication, monitoring, triage, diagnosis, decision support, and care-related services.	Public evaluations of healthcare AI become more conditional when respondents are prompted to consider people with diverse access and support needs.
Accessibility in healthcare AI is often framed through usability, digital access, and inclusion, but less is known about how lay respondents connect accessibility with acceptable care.	Respondents linked accessibility not only to usability, but also to intelligibility, communicative fit, accountability, fairness, and access to human recourse.
Trust, technology acceptance, age, and digital confidence are known to shape attitudes toward healthcare technologies.	In the selected regression models, technological acceptance and age were the clearest correlates.
Human oversight and accountability are widely recognized as important for responsible healthcare AI.	The qualitative findings show that respondents viewed human fallback, relational care, and safeguards against misunderstanding or misuse as conditions for acceptable AI-mediated healthcare.

affected by how AI use is framed and explained [15]. Public perceptions of healthcare AI reveal both ethical concerns and opportunities for patient-centered care, particularly where AI may affect comfort, trust, shared decision-making, and the perceived human quality of care [16]. Patients have been shown to identify both potential benefits and risks of AI in healthcare, including concerns about human control, communication, responsibility, and the appropriate boundaries of automated systems [17].

Accessibility becomes especially salient when AI is considered in relation to people whose access to healthcare is shaped by disability, age-related change, chronic illness, cognitive or sensory difference, linguistic barriers, or other conditions that affect how care is accessed, understood, and navigated. Despite increased attention, a noticeable gap remains for detailed exploration of the socio-technical and inclusivity-related implications of AI in healthcare, especially from the perspective of diverse access and support needs. This gap is important for healthcare informatics because anticipated implementation problems may reveal where AI systems are likely to fail users, where trust is fragile, and which forms of recourse and governance are expected before deployment becomes acceptable.

We use diverse access and support needs as an analytic umbrella for individuals whose engagement with healthcare and digital systems may involve physical, cognitive, sensory, communicative, linguistic, cultural, age-related, situational, or other factors that shape access, understanding, and support.

We ask:

- How are accessibility-related evaluations of AI-mediated healthcare patterned when respondents are prompted to consider people with diverse access and support needs?
- Which demographic and attitudinal factors are associated with more supportive or more concerned evaluations in this context?
- How do respondents who endorsed concern describe risks and conditions of acceptable AI-mediated care when considering people with diverse access and support needs?

2. Methods

We designed a mixed-methods survey (see [Supplementary Material 1](#) for structure) to examine public perceptions of AI in healthcare and well-being. Data were collected between May 12 and May 23, 2023. We screened responses for low-effort response patterns, including straightlining and near-zero within-person variance, and conducted sensitivity

checks excluding flagged cases. A total of 1,151 respondents completed the survey.

2.1. Mixed-methods survey

The survey contained 167 questions and typically took approximately 15–20 min to complete. It included numerical, single-choice, multiple-choice, Likert-type, and open-text questions. The survey was implemented using LimeSurvey version 6.0 (LimeSurvey GmbH) on a GDPR-compliant server in Germany and was offered in Finnish and English. Data were collected through a third-party survey provider with access to a large pool of Finnish consumers.

The survey consisted of three parts. Part one focused on background information, including gender, personality, health status, and technological competence.

Part two examined opinions about applications of AI in healthcare. Respondents were asked to evaluate eight healthcare AI use cases spanning self-care, monitoring, robotic surgery, mental health, bio-electronic monitoring, and robotic caregiving. Five cases were randomly selected for each participant to reduce cumulative and order effects. Respondents were asked to consider use-case-specific trust and acceptance judgements. Results focusing on part 2 of the survey have been reported in Kauttonen et al. [18].

The present article focuses on part 3 of the survey, which examined overall evaluations of AI in healthcare in relation to people with diverse access and support needs and on the open-ended elaborations of concern. We focus on six accessibility-related items:

- People of all types of abilities should be able to easily use AI
- AI should be usable and accessible for everyone regardless of abilities
- Accessible AI increases opportunities for both health and well-being practitioners and patients with disabilities
- AI should be able to process input and offer output for people of diverse abilities
- AI can support people with diverse needs
- I feel concerned about people with special needs interacting with AI

The qualitative analysis examines responses to the optional open-ended question that followed the concern item: “What concerns do you have about people with special needs interacting with AI?” This open-text field appeared to respondents who answered agree or strongly agree to the concern item.

2.2. Ethics and responsible research

This study was conducted in accordance with GDPR and the Finnish National Board on Research Integrity guidelines for research with human participants [19]. Under these guidelines, prior formal ethics committee review was not required. Participants received study information before participation and provided informed consent through the provider’s standard procedures. Participation was voluntary, and potentially sensitive items were optional. The provider managed recruitment, compensation, consent administration, and identity data, while the researchers received only de-identified responses and had no access to names, contact details, panel identifiers, or other direct identifiers. Each respondent received a small monetary reward through the provider’s standard procedures. Results are reported only in aggregate form or through anonymized translated excerpts.

2.3. Quantitative analysis

Quantitative analyses examined associations between demographic and attitudinal predictors and four accessibility-related evaluations of AI in healthcare. Analyses were conducted in R using brms/Stan for Bayesian regression and lavaan for confirmatory factor analysis. Correlation analyses used Kendall’s tau with Bayes factors. Three outcomes

were ordinal Likert-scale responses: (1) Accessible AI increases opportunities for both health and well-being practitioners and patients with disabilities; (2) AI can support people with diverse needs; and (3) I feel concerned about people with special needs interacting with AI. These were modeled using cumulative probit regression for ordered categorical responses. A fourth outcome was a CFA-derived factor score representing broad support for usability and accessibility across ability differences.

The primary predictors were age, gender, education level, self-reported health status, health service usage, technological acceptance, and AI knowledge. Age was modeled with linear and quadratic terms, and technological acceptance and AI knowledge were standardized. A categorical self-described technology-adoption variable was included in selected exploratory models. For each outcome, 117 candidate models were compared using leave-one-out cross-validation alongside substantive interpretability.

Full factor loadings, reliability estimates, model details, specifications, priors, convergence checks, coefficients, contrasts, and predictive summaries are provided in [Supplementary Material 3](#).

2.4. Qualitative analysis

The final qualitative dataset consisted of 117 elaborated responses. These were translated from Finnish into English. Among respondents who entered any open-text response, 61.2% were women and 38.8% were men. The median age was 54 years (IQR 38.5–66). The most common education categories in this response-providing subgroup were high school or vocational certificate (43.9%) and undergraduate degree (30.9%).

Our analysis followed codebook thematic analysis from an interpretive qualitative orientation [20]. The analysis proceeded in several stages. First, two members of the research team familiarized themselves with the full set of responses and developed an initial codebook through close reading and discussion. Consensus meetings involved all four researchers and took place after initial familiarization, after each coding round, and at the end of the analysis phase. Second, the same two members of the research team applied and discussed the initial codebook collaboratively, refining code names, definitions, and boundaries as analysis progressed (see [Supplementary Material 2](#)). Third, the same two members of the research team coded the data using a multi-code approach, allowing a single response to receive more than one code where appropriate. Multi-code analysis allowed us to capture the layered character of respondents' concerns.

3. Quantitative results

Quantitative analyses were conducted on $N = 1138$ respondents after excluding two cases with missing data and 11 respondents whose gender category could not be included in the gender-based model term because the subgroup size was too small for stable estimation. In the survey instrument, gender identity options were male, female, nonbinary, and I would rather not say. In the present regression models, only the female and male categories were represented with sufficient cell sizes for stable estimation of the gender term.

The quantitative analyses focus on relationships among four focal accessibility-related evaluations within the analytic sample: broad support for usability and accessibility across ability differences, perceived opportunities associated with accessible AI, perceived supportive potential, and concern about people with special needs interacting with AI. The support-oriented evaluations clustered together more strongly than they aligned with concern. Concern showed only weak association with the broad usability factor, which suggests that it functioned as a partly distinct evaluative dimension.

3.1. Correlation results

All pairwise correlations among the four focal response variables are shown in [Table 2](#). We used Kendall's tau and computed Bayes factors together with lower 5% and upper 95% interval estimates for the correlations.

Five of the six pairs showed evidence of association. The strongest association was between *AI can support people with diverse needs* and *Accessible AI increases opportunities for both health and well-being practitioners and patients with disabilities* with a correlation of 0.557, indicating that these two supportive evaluations tended to co-occur. The concern item showed modest negative associations with the two support-oriented items and little association with the broad usability factor. This pattern suggests that concern operated as a partly distinct evaluative dimension.

3.2. Regression results

We estimated four regression models: one for the broad usability/accessibility factor and three for the ordinal items. Model selection used LOO-CV and substantive interpretability. Across outcomes, 78, 25, 11, and 16 of 117 candidate models fell within one standard error of the top LOO-CV model, indicating substantial model-selection uncertainty. We therefore do not interpret the selected specifications as uniquely supported or as definitive predictor rankings. [Table 3](#) reports exploratory selected-model summaries.

Table 2

Kendall's tau correlations between the four response variables. Here, F = factor variable. Pairs with Bayes Factor at least 100 are marked with *.

Response 1	Response 2	R	BF	5%	95%
AI should be usable for all (F)	Accessible AI increases opportunities for both health and well-being practitioners and patients with disabilities	0.381*	4.875 × 10 ⁷⁹	0.341	0.418
AI should be usable for all (F)	AI can support people with diverse needs	0.387*	2.673 × 10 ⁸²	0.348	0.425
AI should be usable for all (F)	I feel concerned about people with special needs interacting with AI	0.032	0.143	−0.007	0.070
Accessible AI increases opportunities for both health and well-being practitioners and patients with disabilities	AI can support people with diverse needs	0.557*	8.265 × 10 ¹⁷¹	0.516	0.593
Accessible AI increases opportunities for both health and well-being practitioners and patients with disabilities	I feel concerned about people with special needs interacting with AI	−0.124*	1.477 × 10 ⁷	−0.163	−0.086
AI can support people with diverse needs	I feel concerned about people with special needs interacting with AI	−0.132*	2.164 × 10 ⁸	−0.171	−0.094

Table 3
Selected-model summary of regression associations and model-selection uncertainty.

Outcome	Model selection context	Age	Technological acceptance	Other associations	Overall interpretation
AI should be usable for all regardless of abilities	78 of 117 candidate models fell within 1 SE of the top LOO-CV model	Negative quadratic association; endorsement decreased more sharply at older ages	Positive; higher technological acceptance associated with stronger endorsement	Males tended to agree slightly less. Education $\times$ health-status interaction suggested better health was associated with stronger agreement, especially at graduate level; association was negative at bachelor level. Elementary and bachelor-level education showed higher agreement than graduate or high vocational training.	In these selected specifications broad support for universal usability was strongest among more technologically accepting and younger respondents; education-related patterns were more complex and interacted with health status.
Accessible AI increases opportunities for both health and well-being practitioners and patients with disabilities	25 of 117 candidate models fell within 1 SE of the top LOO-CV model	No clear main age effect retained in final interpretation	Positive; clearest predictor of stronger agreement	Graduate-level health $\times$ education interaction suggested better health was associated with stronger agreement among graduates. Education-related estimates varied by level, but intervals were wide and did not support a clear ordered pattern. Respondents who saw themselves as “behind others” in technology adoption were somewhat more likely to agree.	In these selected specifications perceived opportunity was most clearly associated with technological acceptance; education and self-described technology adoption showed weaker and more uncertain additional patterning.
AI can support people with diverse needs	11 of 117 candidate models fell within 1 SE of the top LOO-CV model	Negative linear and quadratic association; older respondents, especially at the highest ages, were less likely to agree	Positive; more technologically accepting respondents were more likely to endorse AI’s supportive potential	Bachelor-level respondents showed the highest support relative to high vocational training, while high-school/vocational respondents showed the lowest support, though intervals did not support clear separation. Better health status was associated with stronger agreement at graduate level.	In these selected specifications support for AI’s supportive role was patterned most clearly by younger age and higher technological acceptance; education-related differences were suggestive but uncertain.
I feel concerned about people with special needs interacting with AI	16 of 117 candidate models fell within 1 SE of the top LOO-CV model	Positive; older respondents reported greater concern	Negative; more technologically accepting respondents reported less concern	Males reported lower concern. AI knowledge showed no clear main association, but its interaction with age suggested higher AI knowledge was associated with lower concern among older respondents.	In these selected specifications concern was patterned most clearly by older age and lower technological acceptance; AI knowledge appeared relevant mainly in interaction with age.

Taken together, the selected regression models suggest that technological acceptance and age were the clearest correlates of accessibility-related evaluations in these specifications. Higher technological acceptance was associated with stronger support-oriented evaluations and lower concern. Age was associated with lower endorsement of AI’s supportive potential and greater concern in the relevant selected models. Other predictors, including education, health status, technology-adoption style, AI knowledge, and interaction terms, appeared more selectively and should be interpreted as exploratory. Because many candidate models showed similar LOO-CV performance, these findings should not be read as a definitive predictor ranking across the full model space.

4. Qualitative results

Among respondents who endorsed the concern item and provided an open-text response, the qualitative analysis identified four concern-elaboration domains: (1) human recourse and relational care, (2) communicative accessibility and misunderstanding, (3) safety, reliability, and override capacity, and (4) privacy, misuse, and unfair classification. These themes describe concern elaborations, not the full range of public views. Table 4 summarizes these domains, their sub-themes, associated code clusters, frequencies of occurrence, and illustrative quotations. Frequencies refer to the proportion of qualitative respondents ($n = 117$) whose responses were coded within each sub-theme. Because multi-label coding was used, totals exceed 100% across categories.

4.1. Human recourse and relational care

Respondents were concerned about the continued presence of human care within AI-mediated healthcare. They emphasized that AI systems could not provide the human connection, empathy, and contextual understanding needed in healthcare settings, especially for people with diverse access and support needs. Further, some respondents viewed AI-mediated care as potentially dehumanizing, linking automation to objectification, neglect, and moral erosion. Respondents expressed that AI-mediated care should preserve access to human support and judgment, as well as the related accountability and responsibility, which should fall on human caregivers.

4.2. Communicative accessibility and misunderstandings

Respondents were also worried about misunderstandings between AI and individuals with diverse access and support needs. On one hand, AI might misunderstand individuals with speech impairments, using simple language, dialect variation, indirect expression, and other forms of non-standard communication. They raised concern that systems may miss implication, ambiguity, non-verbal meaning, or distress. On the other hand, individuals with diverse access and support needs might misunderstand AI. Respondents questioned whether some users would know that they were interacting with AI, understand what the system could and could not do, or appreciate the risks of accepting outputs uncritically. Participants emphasized issues related to training, technological familiarity, and unequal capacity to engage with AI-mediated services.

Table 4

Concern-elaboration domains among respondents who agreed or strongly agreed that they felt concerned about people with special needs interacting with AI and provided an open-text response (n = 117).

Implementation-relevant risk domain	Original theme / subtheme	Codes	Frequency of occurrence	Illustrative quotation
1 Human recourse and relational care	1a. Loss of Human Contact & Isolation	“left alone with AI”, “no social satisfaction”, “isolation risk”, “erodes community”, “mental health”	35.90%	“In all these AI questions, what causes concern is people’s ability to use them. In addition, physical touch and humanity are missing. Will this isolate people into their own compartments??” (P1038_F65)
	1b. Needing a Human on Social & Functional Level	“replacing caregivers”, “need for human”, “physical accessibility”	33.33%	“There must be real human interaction alongside. In some night-time elderly care functions, there might be some monitoring with AI, but society has gone completely mad if this replaces human work.” (P595_F67)
	1c. Lack of Empathy & Emotional Warmth	“AI has no empathy”, “cold responses”, “can’t read feelings”, “no comforting touch”, “doesn’t ‘feel’ concern”	21.37%	“AI may not be able to consider individuality. And what if AI delivers facts but can’t soften the message – like telling a seriously ill person they have little chance of survival, without being able to support the patient.” (P799_F25)
	1d. Dehumanization & Dignity Threat	“objectification”, “treated like data”, “people as guinea-pigs”, “moral erosion”, “AI watching elders die alone”	11.11%	“Individuals with diverse needs should specifically be protected from being at the mercy of any AI, which lacks even the minimal sense of responsibility or empathy that human decision-makers have.” (P1528_F25)
2 Communicative accessibility and misunderstanding	2a. Digital Access & Literacy Gaps	“no computer / smartphone”, “don’t know how to use”, “lack of training”, “language barrier”	29.06%	“I don’t believe everyone will learn to use AI properly.” (P592_F62) “Elderly people don’t really understand technology... so-called disabled people may not understand AI and prefer to talk and connect with a real human being...” (P1518_F24)
	2b. Trust, Misunderstanding & Lack of Awareness	“don’t realise it’s AI”, “over-trust”, “take advice literally”, “can’t judge credibility”, “uncritical acceptance”, “trust”, “misunderstanding”, “risk and limitation awareness”	58.97%	“A person must understand who or what they are interacting with. Between people, this is easy, but for the elderly, mental health patients, or people with developmental disabilities, dealing with technology can be unfamiliar and lead to misunderstandings.” (P383_F28)
	2c. AI Limitations in Understanding Human Communication	“speech impairments not recognised”, “simple language misread”, “body-language ignored”, “dialect confusion”	16.24%	“AI doesn’t always understand questions. For many, even language skills and speech are barriers—not to mention understanding.” (P343_F74) “AI doesn’t understand a person whose communication ability is limited!” (P1642_M59) “I don’t believe a machine can yet be programmed to read between the lines, if the person’s response is not clear... some people exaggerate or underplay what they say...” (P1698_M60)
3 Safety, reliability, and override capacity	3a. Errors & Misinformation	“wrong answers”, “faulty medical advice”, “data misread”, “false positives/negatives”, “hallucinations”	23.93%	“I’ve used various companies’ robots (AI) to ask questions so I wouldn’t have to wait on the phone. Almost every time, I didn’t get an answer—wrong keyword, question not understood, etc...” (P298_M64)
	3b. Unfinished / Poorly Tested Tech	“beta products in care”, “not ready for deployment”, “patients as test subjects”, “lack of validation”	14.53%	“AI should be ready enough at the time of deployment so that people don’t end up as test subjects...” (P249_M75)
	3c. Technical Failures & Emergencies	“system crash mid-task”, “power outage”, “internet down”, “AI malfunctions”, “unpredictable behaviour”, “cannot override AI”	9.40%	“What happens when the internet crashes during treatment? A power outage?” (P1321_F75)
4 Privacy, misuse, and unfair classification	4a. Privacy & Data Security	“sold to insurers”, “profiling”, “GDPR worries”, “cyber security risk”	14.53%	“I’m mainly worried about whether large companies, like insurance companies, will use the data AI collects against individuals with diverse needs...” (P1586_F24)
	4b. Exploitation & Manipulation	“targeting human weaknesses”, “AI persuades costly decisions”, “predatory use of data”, “misuse”	22.22%	“AI reveals the user’s weak points, enabling a third party – a human – to exploit those weaknesses by asking for money or advising them to invest in some online investment scheme.” (P160_M50) “The idea of deceiving AI or people deliberately misusing AI is also something to think about.” (P1375_F35)
	4c. Bias, Discrimination & Fairness	“unfair classification”, “de-prioritised care”, “value of life questions”, “social hierarchy”, “exclusiveness”, “delayed care”	35.04%	“There’s a possibility that individuals with diverse needs could be classified in ways that harm them...” (P1316_M68)
	4d. Need for Human Oversight & Protection	“must safeguard vulnerable”, “AI only supportive”, “human accountability”, “ethical governance”	21.37%	“Individuals with diverse needs should specifically be protected from being at the mercy of any AI, which lacks even the minimal sense of responsibility or empathy that human decision-makers have.” (P1528_F25)

These concerns were especially visible in responses about older adults and people with cognitive impairments, who were imagined as being left behind if healthcare services became organized around systems they could not easily use or interpret.

4.3. Safety, reliability, and override capacity

Respondents questioned whether current AI technologies were sufficiently ready for use where users may be dependent on the system, may have limited capacity to detect errors, or may face serious consequences if something goes wrong. Errors and misinformation worried respondents, including fears of wrong answers, faulty advice, false positives or negatives, and misread data. Respondents treated these risks as especially serious where the user might not be able to recognize, interpret, or challenge an incorrect output independently. Across these concerns, quality control, human oversight, and mechanisms for intervention when AI-mediated care breaks down were seen as crucial.

4.4. Privacy, misuse, and unfair classification

Concerns related to AI systems handling health-related or disability-related information were worrisome to respondents. Privacy and data security, as well as the inappropriate use of data through profiling, surveillance, or secondary use by institutions such as insurers or other actors were flagged by respondents. Relatedly, respondents worried that AI systems could classify people in harmful ways, deprioritize them, or reproduce exclusion through automated judgments. Some responses extended this concern into a broader critique of care organized around automation, suggesting that reliance on AI could signal reduced social commitment to people with diverse access and support needs.

5. Discussion

The findings suggest that implementation quality depends partly on whether systems can be understood and used safely by people whose communication, cognition, literacy, sensory access, language, or confidence may differ from dominant design assumptions. The qualitative analysis clarifies how respondents framed accessibility as communicative and organizational fit, not only technical usability. Respondents repeatedly described risks of mutual misunderstanding. This aligns with earlier work showing that AI accessibility is shaped by technical availability, digital confidence, communicative fit, and practical conditions of use [4–6].

The findings also align with accessibility standards such as WCAG 2.2, ISO 9241–171, and EN 301 549, which foreground perceivable, operable, understandable, and robust interaction, including understandable communication, flexible input/output, and usable alternatives [21–23]. Yet healthcare AI also raises issues beyond interface-level accessibility, including clinical responsibility, human fallback, safety, institutional accountability, and misclassification. Predeployment review should therefore include subgroup-sensitive assessment of speech, language, literacy, intelligibility, sensory access, and user understanding of system limits, alongside training, onboarding, plain-language communication, multimodal interaction, and accessible non-AI alternatives.

The dominant qualitative theme concerned human recourse and relational care. Our findings align with patient and public attitude research, where support for healthcare AI often weakens when systems seem to displace human judgment, reduce the human presence of care, or assume roles that should remain under accountable human oversight [13,16,17,24–26]. Respondents in our data described conditions for acceptable use, including bounded system roles, access to human professionals, review of consequential outputs, override capacity, and accountability when systems fail. This is consistent with WHO guidance emphasizing human autonomy, safety, transparency, responsibility, accountability, inclusiveness, and responsiveness to human needs [27],

and with regulatory guidance stressing risk management, evaluation, documentation, transparency, and post-deployment monitoring [28]. Human recourse should therefore be understood as part of the accessibility, acceptability, and patient-centeredness of healthcare AI.

Health-equity literature shows that AI systems can reproduce or worsen inequities across the AI lifecycle unless data, model design, deployment, transparency, governance, and community involvement are explicitly addressed [29]. Disability-specific fairness work adds that healthcare AI can create distinct risks when data, communication patterns, embodiment, support arrangements, and care trajectories are poorly represented in algorithmic systems [30]. Respondents' concern elaborations in our data identify possible exclusionary mechanisms, which support fairness evaluation that treats accessibility and disability as more than demographic variables. This suggests that public evaluations may become more conditional when respondents consider users who face additional barriers in digital or AI-mediated care [14,15].

Several limitations should be noted. Findings from a Finnish web-survey sample may not generalize to other healthcare systems or populations. The study examined public evaluations under an access-needs prompt and does not provide direct evidence about the experiences, preferences, or accessibility requirements of people with specific disabilities or support needs. Diverse access and support needs were used as a broad sensitizing prompt. The open-text material came only from respondents who agreed or strongly agreed with the concern item and then provided an elaboration. The open-text subgroup was 61.2% women, so themes related to empathy, relational care, and human contact may partly reflect the composition of respondents. The regression analyses involved substantial model-selection uncertainty and should be interpreted as exploratory selected-model associations. Small cell sizes meant that nonbinary and “rather not say” gender categories could not be modeled separately. Respondents also evaluated five of eight randomly assigned AI use cases before answering the general accessibility items, so their evaluations may have been influenced by the specific use-case set they encountered. The study also did not evaluate actual healthcare AI systems, accessibility compliance, clinical outcomes, or real interactions between AI systems and people with access or support needs.

Future research should directly include people with diverse access and support needs through participatory, co-produced, and subgroup-sensitive designs.

Funding statement

This work was supported by the Research Council of Finland through the project Emotional Experience of Privacy and Ethics in Everyday Pervasive Systems (BUGGED) (decision no. 348391), and by the Ministry of Education and Culture through the project AI Forum: Tekoölyvusteinen digitaalinen murros (AI Forum) (grant no. OKM/236/523/2020). The funders had no involvement in the conducted research beyond financial support.

Data availability

The data that support the findings of this study are not publicly available because they contain potentially sensitive survey material and are subject to contractual and data-protection restrictions.

Declaration of generative AI and AI-assisted technologies in the manuscript preparation process

During the preparation of this work the authors used ChatGPT 5.5 in order to support language revision and drafting. After using this service, the authors reviewed and edited the content as needed and take full responsibility for the content of the published article.

CRediT authorship contribution statement

Joni-Roy Piispanen: Writing – review & editing, Writing – original draft, Visualization, Project administration, Methodology, Investigation, Formal analysis, Data curation, Conceptualization. **Janne Kauttonen:** Writing – review & editing, Writing – original draft, Visualization, Software, Project administration, Methodology, Investigation, Funding acquisition, Formal analysis, Data curation, Conceptualization. **Ari Alamäki:** Writing – review & editing, Writing – original draft, Validation, Supervision, Methodology, Conceptualization. **Rebekah Rousi:** Writing – review & editing, Writing – original draft, Validation, Supervision, Methodology, Funding acquisition, Formal analysis, Conceptualization.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgements

None.

Appendix A. Supplementary data

Supplementary data to this article can be found online at <https://doi.org/10.1016/j.ijmedinf.2026.106625>.

References

- [1] W. Huo, G. Zheng, J. Yan, L. Sun, L. Han, Interacting with medical artificial intelligence: Integrating self-responsibility attribution, human–computer trust, and personality, *Comput. Hum. Behav.* 132 (2022) 107253, <https://doi.org/10.1016/j.chb.2022.107253>.
- [2] K. Liu, D. Tao, The roles of trust, personalization, loss of privacy, and anthropomorphism in public acceptance of smart healthcare services, *Comput. Hum. Behav.* 127 (2022) 107026, <https://doi.org/10.1016/j.chb.2021.107026>.
- [3] D. Giansanti, A. Pirrera, Integrating AI and assistive technologies in healthcare: Insights from a narrative review of reviews, *Healthcare* 13 (5) (2025) 556, <https://doi.org/10.3390/healthcare13050556>.
- [4] J.G.O. Marko, C.D. Neagu, P.B. Anand, Examining inclusivity: the use of AI and diverse populations in health and social care: a systematic review, *BMC Med. Inf. Decis. Making* 25 (2025) 57, <https://doi.org/10.1186/s12911-025-02884-1>.
- [5] E. Goldenthal, J. Park, S.X. Liu, H. Mieczkowski, J.T. Hancock, Not all AI are equal: Exploring the accessibility of AI-mediated communication technology, *Comput. Hum. Behav.* 125 (2021) 106975, <https://doi.org/10.1016/j.chb.2021.106975>.
- [6] M.R. Morris, AI and accessibility, *Commun. ACM* 63 (6) (2020) 35–37, <https://doi.org/10.1145/3356727>.
- [7] N.G. Packin, Disability Discrimination using Artificial Intelligence Systems and Social Scoring: can we Disable Digital Bias? *Journal of International and Comparative Law* 8 (2) (2021) 487–512. <https://heinonline.org/HOL/P?h=hein.journals/jintcl&i=491>.
- [8] M. Whittaker, M. Alper, C. Bennett, S. Hendren, E. Kaziunas, M. Mills, M. Morris, J. Rankin, E. Rogers, M. Salas, S. West, Disability, Bias & AI Report, AI Now Institute, 2019 <https://ainowinstitute.org/publications/disabilitybiasai-2019>.
- [9] O. Almatrafi, A. Johri, H. Lee, A systematic review of AI literacy conceptualization, constructs, and implementation and assessment efforts (2019–2023), *Comput. Educ. Open* 6 (2024) 100173, <https://doi.org/10.1016/j.caeo.2024.100173>.
- [10] D.T.K. Ng, J.K.L. Leung, S.K.W. Chu, M.S. Qiao, Conceptualizing AI literacy: an exploratory review, *Comput. Educ.: Artif. Intell.* 2 (2021) 100041, <https://doi.org/10.1016/j.caeai.2021.100041>.
- [11] S.S. Qiu, L. Zhang, F. You, X. Zhao, Unpacking media channel effects on AI perception: a network analysis of AI information exposure across channels, overload, literacy, and anxiety among chinese users, *Comput. Hum. Behav.* 173 (2025) 108790, <https://doi.org/10.1016/j.chb.2025.108790>.
- [12] C.E. Tseng, S.H. Jung, Y.N. Elglaly, Y. Liu, S. Ludi, Exploration on integrating accessibility into an AI course, in: *In Proceedings of the 53rd ACM Technical Symposium on Computer Science Education — Volume 1*, 2022, pp. 864–870, <https://doi.org/10.1145/3478431.3499399>.
- [13] A.T. Young, D. Amara, A. Bhattacharya, M.L. Wei, Patient and general public attitudes towards clinical artificial intelligence: a mixed methods systematic review, *The Lancet Digital Health* 3 (9) (2021) e599–e611, [https://doi.org/10.1016/S2589-7500\(21\)00132-1](https://doi.org/10.1016/S2589-7500(21)00132-1).
- [14] B. Beets, T.P. Newman, E.L. Howell, L. Bao, S. Yang, Surveying public perceptions of artificial intelligence in health care in the United States: Systematic review, *J. Med. Internet Res.* 25 (2023) e40337, <https://doi.org/10.2196/40337>.
- [15] C. Robertson, A. Woods, K. Bergstrand, J. Findley, C. Balser, M.J. Slepian, Diverse patients' attitudes towards artificial intelligence (AI) in diagnosis, *PLOS Digital Health* 2 (5) (2023) e0000237, <https://doi.org/10.1371/journal.pdig.0000237>.
- [16] K. Witkowski, R.B. Dougherty, S.R. Neely, Public perceptions of artificial intelligence in healthcare: Ethical concerns and opportunities for patient-centered care, *BMC Med. Ethics* 25 (2024) 74, <https://doi.org/10.1186/s12910-024-01066-4>.
- [17] O. Musbahi, L. Syed, P. Le Feuvre, J. Cobb, G. Jones, Public patient views of artificial intelligence in healthcare: a nominal group technique study, *DIGITAL HEALTH* 7 (2021) 20552076211063682, <https://doi.org/10.1177/20552076211063682>.
- [18] J. Kauttonen, R. Rousi, A. Alamäki, Trust and acceptance challenges in the adoption of AI applications in health care: quantitative survey analysis, *J. Med. Internet Res.* 27 (2025) e65567, <https://doi.org/10.2196/65567>.
- [19] Finnish National Board on Research Integrity TENK. (2019). The ethical principles of research with human participants and ethical review in the human sciences in Finland: Finnish National Board on Research Integrity TENK guidelines 2019 (2nd ed.). Publications of the Finnish National Board on Research Integrity TENK, 3/2019.
- [20] V. Braun, V. Clarke, Using thematic analysis in psychology, *Qual. Res. Psychol.* 3 (2) (2006) 77–101, <https://doi.org/10.1191/1478088706qp0630a>.
- [21] World Wide Web Consortium Web Content Accessibility Guidelines (WCAG) 2.2 W3C Recommendation. 2024.
- [22] International Organization for Standardization. (2025). ISO 9241-171:2025: Ergonomics of human-system interaction. Part 171: Software accessibility (2nd ed.). International Organization for Standardization.
- [23] European Telecommunications Standards Institute. (2021). EN 301 549 V3.2.1: Accessibility requirements for ICT products and services. European Telecommunications Standards Institute.
- [24] P. Esmaeilzadeh, T. Mirzaei, S. Dharanikota, Patients' perceptions toward human–artificial intelligence interaction in health care: Experimental study, *J. Med. Internet Res.* 23 (11) (2021) e25856, <https://doi.org/10.2196/25856>.
- [25] A. Katirai, B.A. Yamamoto, A. Kogetsu, K. Kato, Perspectives on artificial intelligence in healthcare from a Patient and Public Involvement Panel in Japan: an exploratory study, *Front. Digital Health* 5 (2023) 1229308, <https://doi.org/10.3389/fdgh.2023.1229308>.
- [26] Tran, V.-T., Riveros, C., & Ravaut, P. (2019). Patients' views of wearable devices and AI in healthcare: Findings from the ComPaRe e-cohort. *npj Digital Medicine*, 2, Article 53. doi: 10.1038/s41746-019-0132-y.
- [27] World Health Organization, Ethics and governance of artificial intelligence for health: WHO guidance. World Health, Organization (2021).
- [28] World Health Organization, Regulatory considerations on artificial intelligence for health, World Health (2023). Organization.
- [29] C.T. Berdahl, L. Baker, S. Mann, O. Osoba, F. Girosi, Strategies to improve the impact of artificial intelligence on health equity: Scoping review, *JMIR AI* 2 (2023) e42936, <https://doi.org/10.2196/42936>.
- [30] Y. Vogt, Disability and algorithmic fairness in healthcare: a narrative review, *Journal of Medical Artificial Intelligence* 8 (2025) 56, <https://doi.org/10.21037/jmai-24-415>.