

Information literacy in the age of generative tools: Development and validation of the AI Information Literacy Scale (AILIS)

Lilach Alon ^{*} , Inbar Levkovich

Faculty of Education and Teaching, Tel-Hai, University of Kiryat Shmona in the Galilee, Kiryat Shmona, Israel

ARTICLE INFO

Keywords:

AI information literacy
GenAI tools
Information behavior
Information practices
Scale development

ABSTRACT

AI information literacy describes people's information practices, namely how they seek, create, process, and evaluate information with GenAI tools across life contexts. This paper presents the development and validation of the AI Information Literacy Scale (AILIS), a self-report measure that explicitly grounds AI literacy in information behavior traditions and spans the information lifecycle. An initial 43 item pool across six conceptual dimensions was refined through expert review to 40 items and five dimensions, then reduced by principal components analysis to a four factor, 39 item scale, administered to 758 adults in Israel. The three broad factors (creation and processing, assessment and critique, seeking and retrieval) showed excellent internal consistency ($\alpha = .92-.96$), and the two item ethics factor showed acceptable reliability ($\alpha = .73$). Participants reported moderate AI information literacy, with higher scores among younger and more educated respondents and among those who used AI tools more frequently. AILIS offers a cross contextual measure that links AI literacy to concrete information practices and supports research and organizational diagnosis in educational and information intensive work contexts.

1. Introduction

Generative artificial intelligence (GenAI) systems have quickly become major intermediaries in contemporary information environments (Sun et al., 2024). People now turn to conversational agents and image generators to search, summarize, translate, and compose text, code, and visuals as part of their everyday work, learning, and leisure (Alon & Malinoff, 2025). From an information science perspective, these systems are not only new tools. They are powerful information infrastructures that shape how information is produced, organized, and made available, and they reconfigure familiar activities of seeking, evaluating, and managing information (Hong, 2025; Huvila & Gorichanaz, 2025).

The core competence at the center of this study is information literacy, understood in information science as a constellation of abilities for locating, processing, interpreting, and critically evaluating information (Trixa & Kaspar, 2024). In an AI-mediated environment, information literacy also involves recognizing how content is generated, assessing transparency and trustworthiness, and distinguishing between information created by humans and information created by machines (Carroll & Borycz, 2024).

In this study, we conceptualize AI information literacy as a specific configuration of information literacy within AI-mediated information practices. AI information literacy refers to the ability to locate, create, process, and critically evaluate information with and through AI tools, while making informed and ethical decisions about their use. It is not limited to knowing basic AI concepts or holding positive or negative attitudes toward AI (Carolus et al., 2023; Koch et al., 2024). Instead, we frame AI information literacy as a capacity that brings together meta-cognitive strategies such as reflection and monitoring, calibrated trust in AI systems, and concrete information practices such as querying, comparing sources, organizing outputs, and detecting errors when using AI tools for learning, work, and everyday problem solving.

Empirical work links higher levels of information literacy to more informed decision making (Dauer et al., 2022), stronger critical thinking (Al-Zou'bi, 2021), and better detection of biased or misleading information (John & Tater, 2025; Levkovich et al., 2024). These outcomes are central concerns in information behavior research, which examines how individuals seek, use, and manage information across everyday life, work, and learning contexts (Case & Given, 2016; Wilson, 1999). When GenAI systems become primary gateways to information, information literacy can no longer be examined only in relation to documents or

^{*} Corresponding author.

E-mail addresses: alonlil@telhai.ac.il (L. Alon), levkovinb@telhai.ac.il (I. Levkovich).

<https://doi.org/10.1016/j.chbah.2026.100254>

Received 10 December 2025; Received in revised form 19 January 2026; Accepted 6 February 2026

Available online 13 February 2026

2949-8821/© 2026 The Authors. Published by Elsevier Inc. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

conventional search engines. Users now engage in iterative dialogues with systems that may hallucinate (Maleki et al., 2024), embed opaque biases (Wei et al., 2025), or blur the origins of their outputs (Williamson & Prybutok, 2024). This raises new questions about how people monitor, correct, and integrate AI generated information into their own information collections and routine information behavior.

Existing measurement tools only partially capture this construct and are rarely grounded explicitly in the information behavior and personal information management (PIM) literature. Many instruments conceptualize AI literacy primarily in terms of technical knowledge, general digital competence, or attitudes toward AI, typically in educational contexts such as teachers, students, or higher education (Hornberger et al., 2025; Jorolan et al., 2025). As a result, they do not systematically map the full information lifecycle of seeking, creating, processing, storing, and critically appraising AI-mediated information across different life contexts, including but not limited to education. Against this backdrop, it is important to situate new measures of AI-related competence in relation to existing scales.

To address this gap from an information behavior perspective, we developed the AI Information Literacy Scale (AILIS) as a self-report instrument that operationalizes AI information literacy through concrete information practices that individuals can perform with GenAI tools. The aim of this paper is to situate AI information literacy within the traditions of information literacy, information behavior, and personal information management, and to present the development, initial validation of AILIS, and initial findings from its use in a large adult sample.

Specifically, we outline the theoretical foundations that connect AI information literacy to established constructs in information behavior, describe the item development and scale construction process, examine the underlying factor structure and reliability of the scale, and present initial patterns in AI information literacy across demographic groups, levels of AI use and perceived knowledge. By offering a measure that centers information practices rather than only general AI literacy or attitudes, we seek to provide researchers and practitioners with a tool for assessing how people navigate AI-mediated information environments and for designing interventions that support critical, ethical, and effective information management in the age of AI.

2. Theoretical background

2.1. Existing AI literacy and related measurement scales

In recent years, instruments that assess people's competencies for interacting with AI systems have grown rapidly. This work spans both conceptual frameworks that define AI literacy (Long & Magerko, 2020) and empirical instruments that operationalize it (Hornberger et al., 2025; Hou et al., 2025). For the purpose of positioning AILIS, it is useful to distinguish between measures that focus on conceptual AI knowledge, broad AI competencies and dispositions, context specific AI use patterns, and information literacy competencies enacted when GenAI mediates seeking, evaluating, and using information. The present study focuses on the last of these, by examining GenAI information literacy, defined here as information literacy self-efficacy in GenAI-mediated information behavior.

Reviews of AI literacy measures describe a fragmented landscape of self-report scales that rest on different definitions and are often tailored to specific populations such as university students, teachers, or professionals (Jorolan et al., 2025; Lintner, 2024; Ng et al., 2021). Across these instruments, AI literacy is commonly framed as a set of competencies that enable individuals to understand what AI is, recognize its applications, evaluate its implications, and use or design AI systems responsibly.

Alongside these instruments, major frameworks increasingly situate AI literacy within broader digital competence and information literacy discussions. DigComp 3.0, for example, defines digital competence as critical and responsible engagement with digital technologies and

explicitly includes information and data literacy, while integrating AI competence across competence areas (Cosgrove & Cachia, 2025). The European Commission and OECD's AILit Framework for primary and secondary education (OECD, 2025) similarly aims to establish a shared understanding of AI literacy through competence descriptions and learning scenarios that emphasize meaningful and ethical use. These efforts reflect a shift toward linking AI literacy with how people search for, evaluate, manage, and use information in AI-mediated environments.

Academic conceptualizations similarly define AI literacy as a broad competence space that includes knowledge, skills, and critical engagement. Long and Magerko (2020), for example, frame AI literacy in terms of competencies that enable people to understand and evaluate AI systems and participate meaningfully in AI-enabled settings. Ng et al. (2021) synthesize AI literacy as covering understanding, use, evaluation, and ethical issues. These conceptualizations are intentionally broad, spanning outcomes from conceptual knowledge to civic and ethical engagement. AILIS is positioned within this discourse as an instrument that operationalizes the information literacy component of AI literacy, focusing on self-efficacy for GenAI-mediated information behavior across the information life cycle. It is therefore complementary to broader AI literacy measures and does not aim to measure conceptual AI knowledge or AI system design competencies.

Against this conceptual background, we review existing instruments by grouping them according to the primary construct they operationalize. One group of work builds on the broader digital competence and digital self-efficacy literature. Ulfert-Blank and Schmidt (2022) developed a Digital Self-Efficacy scale to assess adults' confidence in using information and communication technologies (ICTs) across domains such as information and data literacy, communication and collaboration, digital content creation, safety, and problem solving. Although this scale explicitly includes information and data literacy, it targets general ICT competences and beliefs and does not address concrete practices of prompting, questioning, or evaluating outputs from GenAI tools.

A second group of studies introduces AI specific literacy scales that cover knowledge, attitudes, and usage in a broad sense. For example, the Meta AI Literacy Scale (MAILS) developed by Carolus et al. (2023) includes dimensions such as using and applying AI, understanding AI, detecting AI generated content, ethical reflection, and creating AI. Koch et al. (2024) provide an overview of existing AI literacy scales and report confirmatory and exploratory factor analysis that reveal recurrent latent dimensions such as use of AI applications, design or programming of AI systems, ethical and critical reflection, and detection of AI versus human content. These instruments conceptualize AI literacy at the level of systems, algorithms, and socio-technical implications and only indirectly address detailed information practices involved in working with GenAI as an information source.

Other scales focus on particular aspects of people's interaction with GenAI. Hou et al. (2025) developed a self-report scale that measures university students' reliance behaviors on GenAI during problem based learning. Their four factors, reflective use, cautious use, thoughtless use, and collaborative use, describe patterns of dependence on AI suggestions across stages of problem identification, information exploration, and solution generation. Jorolan et al.'s (2025) MAD AI scale focuses on teachers' moral ascendancy and dependency in AI integration, with dimensions such as ethical transparency, professional integrity, institutional support, and educator preparedness. Both instruments provide valuable insights into how people position themselves toward AI tools, but they do not systematically map the concrete information practices that unfold when users interact with these tools.

There is also a growing line of performance based AI literacy tests. Hornberger et al. (2025) developed AILIT-S, a short ten item AI literacy knowledge test derived from a longer validated instrument, designed to provide a time efficient objective measure of conceptual AI knowledge for university students. Such tests capture declarative knowledge about AI concepts and applications and not self-efficacy for information

practices with AI tools. Reviews of studies that use these scales show that AI literacy, whether self-reported or performance based, tends to correlate positively with AI self-efficacy, positive attitudes, and digital competencies (Bergdahl & Sjöberg, 2025; Bewersdorff et al., 2025), and negatively with AI anxiety (Cengiz & Peker, 2025; Li et al., 2025). The focus remains on general knowledge and dispositions rather than on concrete information practices in GenAI environments.

In parallel, there is a well-established line of work on information literacy self-efficacy within information science. Serap Kurbanoglu et al. (2006) developed the Information Literacy Self-Efficacy Scale to capture individuals' beliefs in their ability to define information needs, initiate search strategies, locate and access resources, assess and comprehend information, interpret and synthesize it, communicate findings, and evaluate both product and process. Subsequent adaptations of this scale for different contexts, such as De Meulemeester et al.'s (2018) version for medical students, demonstrate how information literacy self-efficacy can be tailored to disciplinary settings while remaining anchored in classical information tasks, for example using bibliographic databases or evidence based medicine resources.

Other validation studies confirm both the robustness and contextual sensitivity of information literacy self-efficacy. Keshavarz et al. (2017) validated a Persian version of the scale with postgraduate students and used factor analysis and structural equation modeling to establish a multidimensional model that differs from the original factor structure but retains high reliability. Naveed and Mahmood (2022) used the scale in a study of business students and found that students reported higher self-efficacy for basic than for advanced information literacy skills. They also showed that age, program of study, computer proficiency, and English proficiency were significant correlates of information literacy self-efficacy. These studies position information literacy self-efficacy as a key construct for understanding how confident people feel in searching, evaluating, and using information in digital environments, and they show that it can be meaningfully related to demographic and contextual variables.

Despite these advances, existing information literacy self-efficacy scales were developed for traditional or web-based information environments and focus on tasks such as using library databases, evaluating websites, or preparing bibliographies. They do not address the distinctive information practices that arise when GenAI tools themselves become primary information intermediaries. Examples of such behaviors include formulating and iteratively refining prompts, negotiating the boundary between human generated and AI generated content, detecting hallucinations and fabrications, managing multilingual interactions, and making ethical decisions about data disclosure and attribution in conversational interfaces (Alon & Malinoff, 2025).

Overall, current AI literacy measures emphasize knowledge, attitudes, or broad usage of AI systems (Carolus et al., 2023; Hou et al., 2025; Jorolan et al., 2025), while information literacy self-efficacy scales emphasize information environments that do not rely on AI as an intermediary (De Meulemeester et al., 2018; Keshavarz et al., 2017; Serap Kurbanoglu et al., 2006). None of these instruments directly operationalize information literacy within GenAI environments as a multidimensional construct grounded in information science.

The AI Information Literacy Scale (AILIS) developed in this study addresses this gap by operationalizing GenAI information literacy as information literacy self-efficacy in GenAI-mediated information behavior. AILIS assesses perceived ability across four information-related competency domains: seeking and retrieval, creation and processing, assessment and critique, and ethical awareness in GenAI use. The scale is designed for everyday, educational, and professional contexts and includes multilingual and metacognitive aspects of working with GenAI. The next section elaborates the conceptual dimensions that structure the scale and link it to the information literacy and information behavior literature.

2.2. Dimensions of the AI information literacy scale (AILIS)

Building on the definition of information literacy outlined in the Introduction (Carroll & Borycz, 2024; Trixa & Kaspar, 2024), we conceptualized AI information literacy as information literacy self-efficacy in GenAI-mediated information behavior. When GenAI tools act as information intermediaries, core information literacy goals remain stable, including recognizing an information need, locating information, evaluating it, and using it to create knowledge and participate ethically (Johnson et al., 2025). What changes is how these goals are enacted. Users must translate needs into prompts, interpret outputs that can be fluent but unverifiable, and manage attribution and provenance in conversational interaction (Liu & Du, 2025; Lv et al., 2025).

To capture these GenAI-specific demands while remaining anchored in information literacy and personal information management scholarship, we organized AILIS around six dimensions that reflect recurring clusters of information practices across the GenAI information lifecycle: creation, processing, assessment and critique, seeking and retrieval, ethics, and identifying the information source. Together, these dimensions specify information-related competencies in GenAI use and distinguish the construct measured by AILIS from broader AI literacy measures that emphasize conceptual AI knowledge or technical AI skills.

The first dimension, *information creation*, refers to using GenAI tools to produce texts, code, images, and other content, and to adapt outputs to specific tasks, audiences, and languages. This dimension aligns with the use and creation aspects of information literacy and with digital content creation competencies (Vuorikari et al., 2022, 2025). These perspectives emphasize the ability to generate, revise, and communicate information products for particular purposes (Ooi et al., 2025) and contexts (Cai & Tian, 2025). In GenAI-mediated environments, creation involves iterative prompting and refinement in which users provide constraints, examples, and feedback to progressively shape the artifact (Shen et al., 2025). Through this interaction, users specify genre, tone, level of detail, and format, and they make decisions about what to keep, remove, or rewrite as the output evolves (Celik et al., 2025; Sigot & Tassoti, 2025). Creation is therefore closely tied to authorship and responsibility, because users remain accountable for aligning the final product with the communicative goal, the intended audience, and relevant norms in the setting where it will be used (Draxler et al., 2024). Co-authoring also supports downstream information practices, since the resulting texts, images, or code are often saved, reused, repurposed, and integrated into personal information spaces and work systems as semi-stable knowledge objects rather than as disposable drafts (Geroimenko, 2026; Hutt et al., 2024).

The second dimension, *information processing*, focuses on how users transform, integrate, and organize AI-generated information. It includes summarizing long outputs, comparing answers from different tools or sources, synthesizing information into coherent products, and adapting content for different formats and audiences (Dinsmore & Fryer, 2026). In GenAI-mediated work, processing also involves judging what should be retained, edited, or discarded, and documenting how AI-produced content was modified before it is reused or shared. Because GenAI can generate plausible but uneven content, users often need to restructure outputs into traceable notes, outlines, tables, or drafts that support later verification and revision (Alon et al., 2026). This corresponds to the information use phase in information behavior models, where individuals integrate new information into existing knowledge structures (Case & Given, 2016). It also resonates with information literacy self-efficacy studies that distinguish basic location skills from higher order abilities to interpret and synthesize information (Lee et al., 2025; Naveed & Mahmood, 2022).

The third dimension, *information assessment and critique*, captures users' capacity to evaluate the quality, accuracy, and bias of AI generated information (Li & Yuan, 2026). It includes critical editing of outputs, identifying missing evidence or inconsistencies, detecting logical flaws and hallucinations, checking timeliness, and recognizing biased

framing. This dimension builds on the evaluation component of information literacy and on critical information literacy perspectives that stress the need to interrogate sources, power relations, and representational practices (Alon, 2025; Elmborg, 2006). In AI-mediated environments, assessment and critique should also address risks that are specific to generative models, such as fabricated references and answers that sound plausible but are incorrect (Bewersdorff et al., 2025; Hornberger et al., 2025; Koch et al., 2024). Recent studies in mental health contexts show that large language models (LLMs) can provide diagnostic and treatment suggestions that sometimes resemble expert opinions but also differ from clinicians and the public in important ways, underscoring the need to carefully examine AI generated recommendations in sensitive settings (Elyoseph & Levkovich, 2024; Levkovich, 2025).

The fourth dimension, *information seeking and retrieval*, captures the ability to clarify an information need, formulate and refine prompts, and decide when and how to use AI tools for searching (Cai & Tian, 2025; Hong, 2025). This dimension is consistent with models of the information search process that stress iterative query formulation and monitoring of relevance (Kuhlthau et al., 2008; Ravuri & Mardis, 2025; Savolainen, 2008). These models describe seeking as an evolving activity in which users move from an initially vague need to progressively more focused queries and adjust strategy as understanding develops. It also aligns with information literacy perspectives that emphasize purposeful searching, strategic selection of tools, and informed navigation of digital information environments (Vuorikari et al., 2025). From a PIM perspective, it reflects how users initiate the flow of new information into their personal workspace by choosing appropriate tools and strategies (Jones et al., 2015). In the context of GenAI, seeking and retrieval also include strategic use of multiple languages and tools to broaden or narrow the search space and to request sources, examples, or clarifications from the system (Alon & Krtalić, 2024).

The fifth dimension, *information ethics*, addresses how users manage issues of transparency, privacy, and attribution when working with AI tools (Radanliev, 2025; Singhal et al., 2024). It captures the ethical judgment involved in deciding what information to disclose to a system, how to protect sensitive or personal data (Zhu et al., 2025), and how to communicate AI involvement in ways that align with institutional expectations and professional norms (Rezaei et al., 2025). Items in this dimension therefore reflect practices such as indicating GenAI use when required, avoiding the entry of confidential content into prompts, and distinguishing one's own contribution from AI-generated content in written, visual, or computational products. This dimension links AI information literacy to information ethics and to AI literacy frameworks that treat responsible use and ethical reflection as core competencies for engaging with AI systems (Carolus et al., 2023; Long & Magerko, 2020).

The sixth dimension, *identifying the information source*, captures users' capacity to recognize, verify, and document where the information presented in AI generated outputs comes from (Shin et al., 2025). It includes distinguishing between AI generated and human authored content (Boutadjine et al., 2025), checking whether references and citations correspond to real sources and keeping track of which resources were used in producing a given answer (Adel & Alani, 2025). In GenAI settings, this competency matters because answers can sound confident even when it is not clear where the information came from. Users therefore need to check whether sources are real and to document what the answer is based on (Messeri & Crockett, 2024). Recent work on AI-mediated information practices extends these ideas to contexts in which content may be composed, remixed, or fabricated by algorithms, requiring users to ask how the system constructed its output and which underlying sources, if any, were used (Carroll & Borycz, 2024; Koch et al., 2024).

Table 1 summarizes the conceptual dimensions of the AI Information Literacy Scale.

Table 1
Conceptual dimensions of the AI Information Literacy Scale (AILIS).

Dimension	Definition	Key references
<i>Information creation</i>	Using GenAI to produce and adapt content such as texts, code, images, and learning materials, including specifying style, length, language, and audience.	Hutt et al., 2024; Geroimenko, 2026; Vuorikari et al., 2025
<i>Information processing</i>	Transforming, integrating, and organizing AI generated information by summarizing, comparing, synthesizing, and adapting outputs for different formats and audiences.	Case & Given, 2016; De Meulemeester et al., 2018; Naveed & Mahmood, 2022
<i>Information assessment and critique</i>	Evaluating AI generated information for coherence, accuracy, timeliness, completeness, hallucinations, and biased or problematic representations.	Bewersdorff et al., 2025; Hornberger et al., 2025
<i>Information seeking and retrieval</i>	Clarifying information needs and formulating, refining, and varying prompts to locate relevant information with GenAI, including choosing between AI and other sources.	Kuhlthau, 1991; Jones et al., 2015; Savolainen, 2008
<i>Information ethics</i>	Managing transparency, privacy, and attribution when working with AI, including indicating AI use, avoiding sensitive data, and distinguishing personal from AI contributions.	Carolus et al., 2023; Long & Magerko, 2020
<i>Identifying the information source</i>	Recognizing, checking, and documenting where information in AI outputs comes from, including distinguishing AI generated from human generated content and tracking provenance.	Carroll & Borycz, 2024; Koch et al., 2024

2.3. Research questions

Building on the conceptualization of AI information literacy and the proposed dimensional structure, the present study focuses on developing and validating AILIS and examining its associations with basic demographic characteristics and perceived AI knowledge and use. Specifically, we address the following research questions.

- RQ1.** What factor structure characterizes AI information literacy, and to what extent do the resulting components of the AI Information Literacy Scale demonstrate acceptable reliability?
- RQ2.** How do AI information literacy components vary across demographic groups (age, gender, and education), and across levels of AI use and perceived AI knowledge in this sample?

By examining these questions, the study aims to provide a robust measurement tool for mapping adults' AI information literacy in diverse settings. A validated AILIS can support needs assessment and evaluation in workplaces and organizations that integrate AI tools, in higher education and professional development programs that seek to foster critical and effective use of AI, and in large-scale surveys that monitor inequalities in access to and use of AI-mediated information across everyday life contexts.

3. Methods

3.1. Sample and procedure

The sample in this study included 758 Israeli adults recruited using

purposive sampling designed to include both male and female participants across a broad age range. Inclusion criteria required that participants be at least 18 years old, reside in Israel, and be able to read Hebrew well enough to complete the questionnaire.

The sample included 394 male (52%) and 364 female (48%), with ages ranging from 18 to 80 years ($M = 43.97$, $SD = 14.24$). Most participants (94.2%) reported Hebrew as their native language. In terms of education, 176 reported having a high-school diploma (23.2%), 163 a professional diploma (21.5%), 260 a bachelor's degree (34.3%), and 159 a master's degree or above (21%).

To characterize experience with AI tools, participants reported the number of different GenAI tools they currently use ($M = 1.51$, $SD = 1.08$). They also rated how frequently they use AI tools on a 1 to 5 scale, with higher scores indicating more frequent use. The mean score was 3.22 ($SD = 1.16$), suggesting moderate use of AI tools on average. All participants completed the questionnaire online in Hebrew and provided informed consent prior to participation.

The survey was administered online via a questionnaire hosted on Google Forms. Data were collected in October and November 2025. Participants accessed the survey through a link distributed via social networks and mailing lists, and the full questionnaire, including AILIS items and background questions, took approximately 10 min to complete. The study was conducted in accordance with institutional and national ethical guidelines for research with human participants. The research protocol was reviewed and approved by the university ethics committee. Participation was voluntary, and informed consent was obtained online prior to starting the survey. Respondents were informed that their answers would be used for research purposes only, that they could discontinue participation at any time, and that no identifying information would be stored with their responses (Elyoseph et al., 2024; Hadar-Shoval et al., 2024). Data were analyzed and reported in aggregate form to protect participants' privacy.

3.2. Scale development

AILIS was developed as a self-report instrument to measure individuals perceived ability to carry out information-related activities with GenAI tools. Item development was guided by the information literacy and information behavior literature, with an emphasis on practices of seeking, processing, creating, and evaluating information in AI-mediated environments.

In the first stage, we constructed an initial pool of 43 statements based on the review of existing tools on information literacy, as described in the theoretical background section. Items were constructed to represent six conceptual dimensions: information creation (9 statements); information processing (9 statements); information assessment and critique (9 statements); information seeking and retrieval (9 statements); information ethics (4 statements); and identifying the information source (3 statements). Each statement described a concrete information practice that could realistically be performed with GenAI tools. Respondents rated how confident they were in their ability to perform each action on a five-point scale from 1 (not at all) to 5 (to a very great extent).

In the second stage, the 43-item version was submitted to six experts in AI literacy and survey development, including two psychologists and three education researchers, all of whom had prior experience in tool development and in working with AI-based systems. The experts were asked to evaluate each statement for clarity and relevance and to comment on the coverage and distinctiveness of the five dimensions. Five of the six reviewers recommended omitting the "identifying the information source" dimension, arguing that its behavioral indicators are likely to become less distinctive as AI detection tools and institutional practices evolve. Consequently, three statements in this dimension were removed from the scale. Other statements were rephrased to better match participants' language (for example, replacing "prompt" with "instructions"). Three additional statements were omitted because

they were repetitive (one in information creation, one in information seeking and retrieval). In addition, two new statements were added: "Create code or an algorithm using an AI tool" (in the information creation dimension) and "Detect factual errors or hallucinations and ask the tool to correct them" (in the information assessment and critique dimension). Finally, one statement was reassigned from the information assessment dimension to the information processing dimension.

The resulting version of AILIS comprised 40 statements organized into five dimensions following expert content review (see Appendix A; pre-factor-analysis item pool): information creation (9 statements); information processing (10 statements); information assessment and critique (9 statements); information seeking and retrieval (8 statements); and information ethics (4 statements). In this version, the information creation dimension focuses on generating and adapting content with AI; the information processing dimension covers summarizing, comparing, and integrating AI outputs; the information assessment and critique dimension concerns checking accuracy, coherence, timeliness, and potential bias in AI-generated information; the information seeking and retrieval dimension includes items about clarifying information needs and formulating effective prompts; and the information ethics dimension addresses transparency, responsible use, and protection of sensitive information. All items were rated on the same five-point confidence scale from 1 (not at all) to 5 (to a very great extent).

At the beginning of the questionnaire, participants received the following instructions: "This questionnaire examines information literacy in AI tools, defined as the ability to locate, process, interpret, and critically evaluate information while using GenAI tools in everyday, educational, and professional contexts. Read each statement and indicate how confident you are in your ability to perform the action described without help from another person. Please rate your confidence from 1 (not at all) to 5 (to a great extent)".

The AILIS items were originally developed in Hebrew to ensure linguistic and cultural appropriateness for Israeli adults. For the purposes of publication and international comparison, we translated the scale into English using a forward and back translation procedure. First, the Hebrew items were translated into English by the first author, a bilingual researcher with expertise in information behavior and AI-mediated practices. A second bilingual colleague independently reviewed the English version, suggesting alternative phrasings where necessary to preserve the nuance of the original items. Next, another bilingual researcher, who had not seen the original Hebrew items, back translated the English version into Hebrew. The back translation was compared with the original Hebrew wording, and minor discrepancies were discussed until consensus was reached on an English version that retained the intended meaning and difficulty level of each item. All data reported in this paper were collected using the Hebrew version of AILIS; the English wording is provided here to facilitate interpretation and potential adaptation in other contexts.

3.3. Data analysis

Data were first screened for accuracy, outliers, and univariate normality. No missing values were detected in the AI information literacy items, so all 758 cases were retained for analysis. We then computed descriptive statistics (means, standard deviations, skewness, and kurtosis) for all items and examined inter-item correlations to assess their suitability for factor analysis.

To address RQ1, we examined the structure and internal consistency of the AI Information Literacy Scale. We first conducted principal components analysis (PCA) with varimax rotation, guided by the theoretical model and the scale development process. An initial five-component solution was specified to correspond to the proposed information literacy dimensions. Items were retained if they loaded .40 or higher on their primary component and did not show substantial cross-loadings. Component scores were then computed by averaging the items

within each dimension or emergent component. Next, we evaluated the reliability of each component and of the total scale by calculating Cronbach's alpha coefficients and corrected item–total correlations, and by inspecting alpha if item deleted to identify poorly performing items.

To address RQ2, we described initial patterns in AI information literacy across participant characteristics. We computed descriptive statistics for the component scores and examined associations between the AILIS components and demographic variables (age, gender, education) and reported AI tool use. Group differences were explored using correlations and group comparison tests, providing an initial picture of how AI information literacy varies across subgroups in this sample.

4. Findings

4.1. Factor structure and internal consistency of AILIS

Building on the five-dimension version of AILIS that emerged from expert review, we used principal component analysis (PCA) as an

exploratory dimension-reduction approach to support initial scale development and to derive a concise and interpretable scoring structure in this sample. One item, “Create code or an algorithm using an AI tool”, loaded strongly on a separate component and showed weak loadings on the other factors. Because it effectively formed a single-item factor and did not contribute to the coherence of the information creation dimension, this item was omitted from the final version of the scale. The remaining 39 items clustered into four components representing information creation and processing, information assessment and critique, information seeking and retrieval, and information ethics. Component loadings, eigenvalues, explained variance, and Cronbach's alpha values are presented in Table 2.

The PCA yielded a refined four-factor structure based on 39 items that explained 68.29% of the total variance. Although the fourth eigenvalue was slightly below 1, this factor was retained because it represented a conceptually coherent information ethics dimension.

The first component represents *information creation and processing*. It includes 16 items that describe generating texts in one or more

Table 2
PCA of the AI information literacy scale (AILIS).

Dimensions	Item no.	Statements	Loadings			
			Factor 1	Factor 2	Factor 3	Factor 4
Information creation & processing (Cronbach's alpha = .963)	1.	Create a text in more than one language using an AI tool.	0.734	0.267	0.212	0.086
	2.	Generate titles or summaries using an AI tool.	0.718	0.267	0.314	0.239
	3.	Translate texts from one language to another.	0.679	0.234	0.329	-0.066
	4.	Instruct the tool how to adapt the style, length of the answer, and language to the task requirements.	0.666	0.343	0.385	0.124
	5.	Compare answers from different tools and identify similarities/differences.	0.648	0.403	0.250	0.140
	6.	Combine information from several AI tools and sources into a coherent text.	0.647	0.329	0.172	0.191
	7.	Adapt an answer generated by an AI tool to a specific audience.	0.640	0.410	0.245	0.307
	8.	Summarize a long answer from an AI tool into main points.	0.633	0.337	0.381	0.286
	9.	Work with answers in different languages and combine them.	0.623	0.385	0.149	0.225
	10.	Ask the tool for examples, exercises, or review questions on a specific topic.	0.608	0.301	0.359	0.284
	11.	Create a text with the help of an AI tool.	0.603	0.227	0.486	-0.02
	12.	Present AI-generated outputs in different formats (for example, text, table, figure).	0.603	0.293	0.278	0.415
	13.	Break down a complex task into steps and instruct the AI tool how to carry it out step by step.	0.586	0.347	0.289	0.383
	14.	Give the tool feedback that guides it to create a more accurate output.	0.573	0.336	0.367	0.185
	15.	Convert a text into a graph or diagram using an AI tool.	0.554	0.317	0.137	0.452
	16.	Create an image or video using an AI tool.	0.528	0.283	0.230	0.021
Information assessment & critique (Cronbach's alpha = .962)	17.	Detect hallucinations and ask the tool to correct them.	0.259	0.796	0.293	0.062
	18.	Identify logical flaws such as overgeneralization, internal contradiction, or misleading phrasing in an AI-generated answer.	0.264	0.778	0.305	0.114
	19.	Recognize factual errors in an AI-generated answer.	0.271	0.776	0.283	0.041
	20.	Identify where evidence, examples, or sources are missing in an AI-generated answer.	0.306	0.757	0.284	0.158
	21.	Evaluate if an AI-generated answer is consistent, clear, and well organized.	0.247	0.745	0.348	0.206
	22.	Critically edit an output generated by an AI tool.	0.362	0.704	0.224	0.236
	23.	Identify biases in the way the tool presents or explains information (for example, gender bias).	0.333	0.675	0.231	0.176
	24.	Check the update date of the information and assess how current it is.	0.259	0.673	0.172	0.109
	25.	Distinguish in an AI-generated answer between fact, opinion, and interpretation.	0.432	0.590	0.306	0.166
	26.	Ask the tool to explain the steps or reasoning that led to its answer.	0.397	0.589	0.287	0.228
	27.	Evaluate the contribution of the AI tool compared with my contribution.	0.319	0.560	0.310	0.423
	28.	Check that credit is correctly attributed to human creators.	0.363	0.553	0.094	0.374
	29.	Verify the originality of a text by comparing it with available sources.	0.503	0.538	0.18	0.154
Information seeking & retrieval (Cronbach's alpha = .923)	30.	Update or refine prompt after receiving an answer to improve the result.	0.275	0.231	0.789	0.136
	31.	Formulate a clear prompt in order to receive appropriate information.	0.211	0.228	0.775	0.207
	32.	Locate information on a topic that interests me with the help of an AI tool.	0.289	0.219	0.745	-0.047
	33.	Decide when it is appropriate to use an AI tool and when to use other tools	0.161	0.339	0.732	0.141
	34.	Clarify for myself what I am looking for before using an AI tool.	0.199	0.229	0.724	0.248
	35.	Ask the AI tool for examples, sources, or citations.	0.346	0.221	0.689	0.077
	36.	Use alternative prompts in order to broaden or narrow my search.	0.262	0.327	0.663	0.060
	37.	Write prompts in different languages to broaden or refine the information.	0.433	0.208	0.474	0.070
Information ethics (Cronbach's alpha = .727)	38.	Avoid entering sensitive information into the tool when creating an output.	0.173	0.459	0.310	0.537
	39.	Indicate my use of AI tools in a transparent way.	0.411	0.355	0.172	0.582
Eigenvalues			22.30	2.25	1.89	0.87
% of Variance			55.75	5.63	4.74	2.17

languages, producing titles and summaries, translating between languages, adapting style and length, presenting outputs in different formats, breaking down complex tasks into steps, giving corrective feedback to the tool, combining outputs from several AI tools, working with answers in different languages, summarizing long answers, and creating images, videos, graphs, or diagrams with AI. Cronbach's alpha for this factor was .963.

The second component reflects *information assessment and critique*. It includes 13 items that describe detecting factual errors, hallucinations, and logical flaws; identifying missing evidence, examples, or sources; evaluating whether an answer is clear, consistent, and well organized; critically editing outputs; identifying bias; checking how current information is; distinguishing between fact, opinion, and interpretation; asking the tool to explain its reasoning; evaluating the contribution of the AI tool compared with one's own contribution; checking that credit is correctly attributed to human creators; and verifying the originality of a text against available sources. Cronbach's alpha for this factor was .962.

The third component corresponds to *information seeking and retrieval*. It includes eight items that describe what one is looking for before using an AI tool, formulating clear prompts, updating or refining prompts after receiving an answer, deciding when it is appropriate to use AI tools versus other tools, locating information on a topic of interest with the help of AI, asking the tool for examples, sources, or citations, using alternative prompts to broaden or narrow a search, and writing prompts in different languages to refine or expand the information space. Cronbach's alpha for this factor was .923.

The fourth component corresponds to *information ethics*. It includes two statements: avoiding entering sensitive information into the tool when creating an output and indicating the use of AI tools in a transparent way. Cronbach's alpha was .727. As noted above, other items that had been initially assigned to the ethics dimension (e.g., credit for human creators and originality checks) were absorbed into the information assessment and critique component.

Overall, these findings provide empirical support for a four-component structure of the AILIS, comprising information creation and processing, information assessment and critique, information seeking and retrieval, and a focused information ethics component. The three broad components showed excellent internal consistency ($\alpha = .92 - .96$), while the two-item ethics component demonstrated acceptable reliability for a brief subscale ($\alpha = .73$).

The final version of the scale, after validation, is provided in [Appendix B](#).

4.2. Patterns in AI information literacy across demographics and AI use

On the four AILIS components, participants reported moderate levels

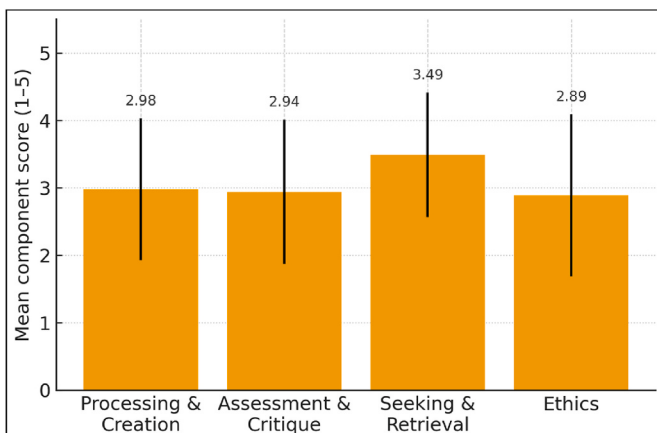


Fig. 1. Mean scores and SD for AILIS components on a 5-point Likert scale ($N = 758$).

of AI information literacy (Fig. 1). Information seeking and retrieval showed the highest mean ($M = 3.49$, $SD = 0.92$), followed by information processing and creation ($M = 2.98$, $SD = 1.05$) and information assessment and critique ($M = 2.94$, $SD = 1.07$). Information ethics had the lowest mean ($M = 2.89$, $SD = 1.20$). For all four components, the full possible range from 1 to 5 was observed, indicating substantial variability in perceived AI information literacy across the sample.

We first examined correlations between the four AILIS components and indicators of AI use (self-reported AI use, knowledge of AI tools, and number of AI tools used; Table 3). All correlations were positive and statistically significant ($p < .01$), indicating that participants who reported higher engagement with AI tools also tended to report higher AI information literacy. Across components, associations with AI use frequency and knowledge of AI tools were in the moderate range: for example, correlations with AI use frequency ranged from $r = .33$ (ethics) to $r = .48$ (processing and creation), and correlations with knowledge of AI tools ranged from $r = .35$ (ethics) to $r = .46$ (processing and creation). Correlations with the number of AI tools used were smaller but still meaningful ($r = .27 - .37$), suggesting that using a broader repertoire of tools is also linked to higher perceived AI information literacy. Age showed the opposite pattern: it was negatively associated with all four components ($r = -.18$ to $-.26$, $p < .01$), indicating that older participants reported somewhat lower AI information literacy and, as reflected in additional correlations, tended to report lower AI use, familiarity with AI tools, and number of tools used.

We next examined differences in AI information literacy across education levels using one-way ANOVA (Table 4). For all four components, the ANOVA tests were significant: processing and creation, $F(3, 754) = 8.70$, $p < .001$; assessment and critique, $F(3, 754) = 6.27$, $p < .001$; seeking and retrieval, $F(3, 754) = 13.02$, $p < .001$; and ethics, $F(3, 754) = 5.57$, $p < .001$.

Bonferroni-corrected pairwise comparisons showed a consistent pattern in which participants in the highest education group (Master's and above) reported significantly higher scores than those in the two lowest education groups (i.e., high school or professional diploma) on all components. For seeking and retrieval, participants with Bachelor's degree also scored higher than those in the two lowest education groups, whereas differences between the two highest education groups were small and non-significant. Overall, AI information literacy tended to increase with education level, with the largest and most consistent gaps observed between the lowest and highest education groups.

Lastly, we examined gender differences, but no significant differences were found for any of the AILIS components.

Overall, these patterns suggest that AI information literacy is systematically related to both demographic characteristics and AI engagement, with higher education and more intensive AI use associated with higher AILIS scores, and older age associated with lower scores across components.

5. Discussion

This study introduced the AI Information Literacy Scale (AILIS) as a self-report measure grounded in information literacy and information behavior. AILIS operationalizes AI information literacy through concrete information practices that individuals perform with GenAI tools. Using data from a large adult sample, we examined the internal structure and reliability of the scale and its associations with demographic characteristics and perceived usage and knowledge of AI tools.

Conceptually, AILIS began from a six-dimension model, which was refined to five dimensions after expert review and to four empirically supported components following principal components analysis. The findings offer an empirically supported and theoretically informed way to describe individuals' AI information literacy, namely how they seek, create, process, and critically evaluate information with GenAI tools across everyday, educational, and professional contexts.

Positioning AILIS in relation to previous tools helps clarify its

Table 3

Pearson correlations among AI information literacy components, AI use, Num of AI tools, AI knowledge and age (N = 758).

Variable	1	2	3	4	5	6	7
1. Processing & creation	-						
2. Assessment & critique	.83**	-					
3. Seeking & retrieval	.76**	.71**	-				
4. Ethics	.74**	.76**	.60**	-			
5. AI Knowledge	.46**	.42**	.44**	.35**	-		
6. Num AI tools	.36**	.30**	.36**	.27**	.35**	-	
7. AI use	.48**	.40**	.45**	.33**	.74**	.37**	-
8. Age	-.26**	-.25**	-.18**	-.22**	-.29**	-.17**	-.30**

Note. All coefficients are Pearson correlations. N = 758 for all variables; AI use frequency = self-reported frequency of using AI tools. ** $p < .01$ (two-tailed).**Table 4**

One-way ANOVAs for AILIS components by education level (N = 758).

Component	Sum of squares (between)	df (between)	Mean square (between)	Sum of squares (within)	df (within)	Mean square (within)	F	p
Processing & creation	27.97	3	9.32	808.41	754	1.07	8.70	<.001
Assessment & critique	20.98	3	6.99	841.50	754	1.12	6.27	<.001
Seeking & retrieval	31.82	3	10.61	614.16	754	0.82	13.02	<.001
Ethics	23.47	3	7.82	1058.63	754	1.40	5.57	<.001

Note. Education was coded into four levels: high school diploma, professional diploma, Bachelor's degree, Master's and above.

contribution. The AI literacy instruments reviewed in the background section largely conceptualize AI literacy as conceptual knowledge about AI technologies, general digital competence, or attitudes toward AI, typically among teachers, students, or other higher education populations (e.g., Carolus et al., 2023; Hornberger et al., 2025; Hou et al., 2025; Jorolan et al., 2025; Koch et al., 2024). These scales are important for understanding what people know about AI and how they feel about integrating it into teaching and learning. However, they are only indirectly connected to information behavior frameworks and they rarely trace the full information lifecycle of seeking, creating, processing, storing, and critically appraising AI-mediated information across different life contexts. By contrast, AILIS starts from information literacy and information behavior traditions and specifies what it means to be information literate in an AI-mediated environment. It shifts the focus from knowledge or perceptions about AI to what people actually do with AI tools when they search for, evaluate, and manage information.

Another way in which AILIS extends previous work concerns the range of information practices it measures. Many existing AI literacy and information literacy scales concentrate on specific aspects, such as understanding basic AI concepts, evaluating the quality of outputs, or reflecting on ethical risks. In contrast, AILIS is designed to span the full information lifecycle with factors that capture how people seek and retrieve information with AI tools, process and create content based on AI outputs, and assess and manage ethical issues related to transparency, bias, and responsible use. This structure suggests that people experience AI information literacy as a combination of broad interaction skills with AI tools and a more focused set of ethical sensitivities.

AILIS is also designed to be cross-contextual rather than tied solely to educational settings. Its items refer to using GenAI for learning, work, and everyday problem solving without presupposing a particular profession or disciplinary background. This makes the scale especially relevant for knowledge workers and organizational contexts, where GenAI is rapidly being embedded in daily workflows (Alon & Malinoff, 2025). Knowledge workers, for whom information and its applications are central to their roles, increasingly use AI to draft, summarize, translate, interpret data, and brainstorm ideas, often in ways that are informal and only partially visible to their organizations (Criado et al., 2025; Haesevoets et al., 2025). In this sense, the scale can serve as a diagnostic tool for organizations that are introducing AI-based systems and wish to understand not only adoption rates but also the quality and distribution of AI-related information practices.

The initial AI information literacy patterns observed in this study underline the importance of such assessment. Differences in AILIS components between education levels, age groups, and levels of perceived AI use and knowledge suggest that AI information literacy is already unevenly distributed among adults. In many workplaces, these differences may translate into disparities in who can effectively leverage AI tools for complex information work, who is more vulnerable to misleading outputs or privacy risks, and who may require additional support (Gkinko & Elbanna, 2022; Schuster & Cotten, 2024).

5.1. Implications

The validated AILIS has implications for research, theory development, and practice. It makes AI information literacy measurable through information behaviors and related competencies. Much of the AI and GenAI literacy literature defines literacy through broad competencies such as conceptual knowledge about AI, general interaction skills, and responsible use, often framed around education or citizenship goals (e.g., Carolus et al., 2023; Koch et al., 2024; Long & Magerko, 2020; Ng et al., 2021).

AILIS sharpens this construct by focusing on the information behaviors that GenAI increasingly mediates in everyday and professional settings, and by distinguishing competencies that are often treated as a single domain. Specifically, it separates seeking and retrieval through GenAI from higher order assessment and critique of outputs and attribution, and from creation and processing practices that involve transforming and integrating information with GenAI. This information literacy specification connects GenAI literacy measurement to information behavior and information literacy scholarship and supports theory development about AI-mediated information behavior.

AILIS also enables more informative empirical analyses because it distinguishes between competencies that do not necessarily develop together. Using separate dimension scores allows researchers to identify uneven profiles, such as strong GenAI-assisted content production alongside weak evaluation and verification skills. It also supports targeted hypothesis testing and intervention evaluation by linking each dimension to relevant outcomes, for example relating assessment and critique to verification behaviors and relating creation and processing to productive GenAI use.

In applied terms, AILIS is suitable for needs assessment and program evaluation in educational and organizational contexts. Dimension-level

scores translate into actionable training priorities. In groups that already use GenAI frequently, the results can inform targeted efforts to strengthen assessment and critique and ethical awareness. In groups that are newly adopting GenAI systems, the results can guide support for seeking and retrieval and basic interaction practices (Figenschou et al., 2024). Retaining an Ethics dimension further signals that responsible GenAI use is integral to information literacy. The current two-item structure is therefore interpreted cautiously and discussed as a limitation, alongside the need for future item development and validation to capture emerging norms and responsibilities associated with GenAI use (Ng et al., 2021; Zhang & Wang, 2025).

5.2. Limitations and future studies

Several limitations of the present study should be acknowledged. The data come from a single national context and focus on Israeli adults, which limits generalizability to other countries and cultural settings. Although the sample included variation in age and education, it does not represent all groups that may be of interest in future research, such as older adults with limited digital access, specific professional communities, or adolescents in school contexts. Participants also differed in their prior knowledge and experience with GenAI tools. As a result, respondents were not necessarily at the same stage of adoption or understanding, which may have shaped how they interpreted and rated some items. Because GenAI tools and use norms continue to evolve, the practices described in the scale and the difficulty they pose to different users may also shift over time.

A further limitation concerns measurement type. AILIS is a self-efficacy instrument and therefore captures perceived capability across GenAI-related information tasks rather than demonstrated performance. Self-efficacy can diverge from what individuals can actually do in practice, and the scale should not be interpreted as direct evidence of objective skill. For these reasons, we interpret AILIS scores as perceptions that are useful for diagnosing perceived strengths and gaps and for evaluating interventions, while not treating them as direct evidence of what participants can do in performance terms.

The single country setting also carries linguistic implications. In Israel, many users routinely shift between Hebrew and English during GenAI use, and Hebrew language prompting or outputs may differ in quality, fluency, and safety filtering compared with English. These language-related differences may influence the information behavior dimensions measured by AILIS, including seeking and retrieval strategies, verification practices, and reliance on AI outputs (Alon & Krtalić, 2025). Future studies should validate AILIS across additional countries and language contexts, examine measurement invariance across language of GenAI use, and test whether relationships between AILIS dimensions and reliance or verification behaviors differ when GenAI is used primarily in English versus other languages.

Methodologically, although the four component scale structure is conceptually coherent and shows strong internal consistency, further validation is needed, for example tests of measurement invariance across demographic groups and contexts, and longitudinal studies. The ethics component, in particular, is currently represented by two items and an eigenvalue below one; future work may consider adding or refining items or modeling ethics as part of a broader critical evaluation factor. In domains such as psychiatry and mental health care, where large language models are already being tested alongside professional judgments for diagnosis, treatment recommendations, and recovery prognosis (Elyoseph & Levkovich, 2024; Levkovich, 2025), adapting and validating AILIS for clinicians and other health professionals will be especially important.

In addition, we used PCA as an exploratory dimension reduction

approach. Future work should therefore verify the factor structure through confirmatory factor analysis in independent samples and examine measurement invariance across key demographic groups and across language of GenAI use. We did not administer a separate digital literacy or digital competence scale in this study and therefore could not directly test discriminant validity between AILIS and general digital literacy. Future studies should administer both types of measures and examine whether AILIS captures GenAI-mediated information competencies beyond general digital skills.

Despite these limitations, AILIS makes a distinctive contribution to emerging work on AI literacy by covering the full information lifecycle and by being applicable across multiple life contexts, including knowledge intensive work. As GenAI tools become part of everyday infrastructures in organizations, educational systems, and daily life, the ability to seek, evaluate, and manage AI-mediated information will become an increasingly important dimension of both individual competence and organizational capacity. AILIS provides a starting point for mapping this dimension, identifying inequalities and gaps, and evaluating interventions designed to foster critical, ethical, and effective use of AI tools in contemporary information environments.

6. Conclusion

This study developed and validated the AI Information Literacy Scale (AILIS) as a multidimensional instrument for assessing information-related competencies in GenAI use. Using expert review and empirical validation in a large adult sample, the results support a coherent factor structure that distinguishes seeking and retrieval, creation and processing, assessment and critique, and ethical awareness as separable components of GenAI information literacy. AILIS provides researchers and practitioners with a practical tool for examining variation in GenAI information literacy, identifying uneven competency profiles, and evaluating targeted educational and organizational initiatives. At the same time, the findings point to priorities for continued scale refinement. The ethics dimension was retained because it reflects a central component of information literacy in GenAI-mediated environments, and its current two-item structure warrants cautious interpretation and further item development. Future research should extend validation across additional populations and contexts and examine how AILIS dimensions relate to behavioral indicators of reliance and verification.

CRedit authorship contribution statement

Lilach Alon: Writing – original draft, Visualization, Validation, Methodology, Formal analysis, Data curation, Conceptualization. **Inbar Levkovich:** Writing – review & editing, Writing – original draft, Visualization, Validation, Methodology, Data curation, Conceptualization.

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

During the preparation of this work, the authors used the GenAI tool ChatGPT (OpenAI) to assist with language editing and improving the clarity of wording in the manuscript. After using this tool, the authors carefully reviewed and edited all content and take full responsibility for the integrity and accuracy of the published article.

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.

Appendix A. AILIS pre-validation version used as input to PCA (40 items)

Dimensions	Item	Statements
Information creation	1.	Create a text in more than one language using an AI tool.
	2.	Generate titles or summaries using an AI tool.
	3.	Translate texts from one language to another with the tool.
	4.	Instruct the tool how to adapt the style, length of the answer, and language to the task requirements.
	5.	Create a text with the help of an AI tool.
	6.	Ask the tool for examples, exercises, or review questions on a specific topic.
	7.	Create code or an algorithm using an AI tool.
	8.	Give the tool feedback that guides it to create a more accurate and improved output.
	9.	Create an image or video using an AI tool.
Information processing	10.	Compare answers from different tools and identify similarities/differences.
	11.	Summarize a long answer from an AI tool into main points.
	12.	Work with answers in different languages and combine them into a single product.
	13.	Present AI-generated outputs in different formats (for example, text, table, figure).
	14.	Break down a complex task into steps and instruct the AI tool how to carry it out step by step.
	15.	Convert a text into a graph or diagram using an AI tool.
	16.	Combine information from several AI tools and sources into a coherent text.
	17.	Adapt an answer generated by an AI tool to a specific audience.
	18.	Distinguish in an AI-generated answer between fact, opinion, and interpretation.
Information assessment & critique	19.	Verify the originality of a text by comparing it with available sources.
	20.	Detect hallucinations and ask the tool to correct them.
	21.	Identify logical flaws such as overgeneralization, internal contradiction, or misleading phrasing in an AI-generated answer.
	22.	Evaluate whether an AI-generated answer is consistent, clear, and well organized.
	23.	Critically edit an output generated by an AI tool.
	24.	Recognize factual errors in an AI-generated answer.
	25.	Identify where evidence, examples, or sources are missing in an AI-generated answer.
	26.	Identify biases in the way the tool presents or explains information (for example, gender bias).
	27.	Check the update date of the information and assess how current it is.
Information seeking & retrieval	28.	Ask the tool to explain the steps or reasoning that led to its answer.
	29.	Update or refine prompt after receiving an answer to improve the result.
	30.	Formulate a clear prompt in order to receive appropriate information.
	31.	Locate information on a topic that interests me with the help of an AI tool.
	32.	Decide when it is appropriate to use an AI tool and when to use other tools
	33.	Clarify for myself what I am looking for before using an AI tool.
	34.	Ask the AI tool for examples, sources, or citations.
	35.	Use alternative prompts in order to broaden or narrow my search.
	36.	Write prompts in different languages to broaden or refine the information.
Information ethics	37.	Avoid entering sensitive information into the tool when creating an output.
	38.	Indicate my use of AI tools in a transparent way.
	39.	Evaluate the contribution of the AI tool compared with my contribution.
	40.	Check that credit is correctly attributed to human creators.

Appendix B. Final validated 39-item AILIS scale

Dimensions	Item	Statements
Information creation & processing	1.	Create a text in more than one language using an AI tool.
	2.	Generate titles or summaries using an AI tool.
	3.	Translate texts from one language to another.
	4.	Instruct the tool how to adapt the style, length of the answer, and language to the task requirements.
	5.	Compare answers from different tools and identify similarities/differences.
	6.	Combine information from several AI tools and sources into a coherent text.
	7.	Adapt an answer generated by an AI tool to a specific audience.
	8.	Summarize a long answer from an AI tool into main points.
	9.	Work with answers in different languages and combine them.
	10.	Ask the tool for examples, exercises, or review questions on a specific topic.
	11.	Create a text with the help of an AI tool.
	12.	Present AI-generated outputs in different formats (for example, text, table, figure).
	13.	Break down a complex task into steps and instruct the AI tool how to carry it out step by step.
	14.	Give the tool feedback that guides it to create a more accurate output.
	15.	Convert a text into a graph or diagram using an AI tool.
	16.	Create an image or video using an AI tool.
Information assessment & critique	17.	Detect hallucinations and ask the tool to correct them.
	18.	Identify logical flaws such as overgeneralization, internal contradiction, or misleading phrasing in an AI-generated answer.
	19.	Recognize factual errors in an AI-generated answer.
	20.	Identify where evidence, examples, or sources are missing in an AI-generated answer.
	21.	Evaluate if an AI-generated answer is consistent, clear, and well organized.
	22.	Critically edit an output generated by an AI tool.
	23.	Identify biases in the way the tool presents or explains information (for example, gender bias).
	24.	Check the update date of the information and assess how current it is.
	25.	Distinguish in an AI-generated answer between fact, opinion, and interpretation.
	26.	Ask the tool to explain the steps or reasoning that led to its answer.
	27.	Evaluate the contribution of the AI tool compared with my contribution.

(continued on next page)

(continued)

Dimensions	Item	Statements
Information seeking & retrieval	28.	Check that credit is correctly attributed to human creators.
	29.	Verify the originality of a text by comparing it with available sources.
	30.	Update or refine prompt after receiving an answer to improve the result.
	31.	Formulate a clear prompt in order to receive appropriate information.
	32.	Locate information on a topic that interests me with the help of an AI tool.
	33.	Decide when it is appropriate to use an AI tool and when to use other tools
	34.	Clarify for myself what I am looking for before using an AI tool.
	35.	Ask the AI tool for examples, sources, or citations.
	36.	Use alternative prompts in order to broaden or narrow my search.
	37.	Write prompts in different languages to broaden or refine the information.
Information ethics	38.	Avoid entering sensitive information into the tool when creating an output.
	39.	Indicate my use of AI tools in a transparent way.

References

Adel, A., & Alani, N. (2025). Can generative AI reliably synthesise literature? Exploring hallucination issues in ChatGPT. *AI & Society*, 40, 6799–6812. <https://doi.org/10.1007/s00146-025-02406-7>

Al-Zou'bi, R. (2021). The impact of media and information literacy on acquiring the critical thinking skill by the educational faculty's students. *Thinking Skills and Creativity*, 39, Article 100782. <https://doi.org/10.1016/j.tsc.2020.100782>

Alon, L. (2025). Information seeking and personal information management behaviors as scaffolding during life transitions: The case of early-career researchers. *Aslib Journal of Information Management*, 77(2), 354–372. <https://doi.org/10.1108/AJIM-01-2023-0027>

Alon, L., & Krtalić, M. (2024). "I try to find comfort in English, but it's hard": Exploring personal information management in multilingual contexts among voluntary migrants. *Journal of Librarianship and Information Science*. , Article 09610006241296896. <https://doi.org/10.1177/09610006241296896>

Alon, L., & Krtalić, M. (2025). "I wish I could use any language as it comes to mind": User experience in digital platforms in the context of multilingual personal information management. *Journal of the Association for Information Science and Technology*, 76(4), 686–702. <https://doi.org/10.1002/asi.24964>

Alon, L., & Malinoff, T. (2025). Understanding barriers to AI use among public sector knowledge workers. *Proceedings of the Association for Information Science and Technology*, 62(1), 1346–1348. <https://doi.org/10.1002/pri.21399>

Alon, L., Malinoff, T., & Levkovich, I. (2026). Information practices and emotion regulation in wartime news consumption. *Journal of Documentation*, 82(1), 154–173. <https://doi.org/10.1108/JD-09-2025-0266>

Bergdahl, N., & Sjöberg, J. (2025). Attitudes, perceptions and AI self-efficacy in K-12 education. *Computers and Education: Artificial Intelligence*, 8, Article 100358. <https://doi.org/10.1016/j.caeai.2024.100358>

Bewersdorff, A., Hornberger, M., Nerdel, C., & Schiff, D. S. (2025). AI advocates and cautious critics: How AI attitudes, AI interest, use of AI, and AI literacy build university students' AI self-efficacy. *Computers and Education: Artificial Intelligence*, 8, Article 100340. <https://doi.org/10.1016/j.caeai.2024.100340>

Boutadjine, A., Harrag, F., & Shaalan, K. (2025). Human vs. machine: A comparative study on the detection of AI-generated content. *ACM Transactions on Asian and Low-Resource Language Information Processing*, 24(2), 1–26. <https://doi.org/10.1145/3708889>

Cai, Y., & Tian, S. (2025). Student translators' web-based vs. GenAI-based information-seeking behavior in translation process: A comparative study. *Education and Information Technologies*, 30, 18997–190259. <https://doi.org/10.1007/s10639-025-13523-7>

Carolus, A., Koch, M. J., Straka, S., Latoschik, M. E., & Wienrich, C. (2023). MAIIS-meta AI literacy scale: Development and testing of an AI literacy questionnaire based on well-founded competency models and psychological change-and meta-competencies. *Computers in Human Behavior: Artificial Humans*, 1(2), Article 100014. <https://doi.org/10.1016/j.chbah.2023.100014>

Carroll, A. J., & Borycz, J. (2024). Integrating large language models and generative artificial intelligence tools into information literacy instruction. *The Journal of Academic Librarianship*, 50(4), Article 102899. <https://doi.org/10.1016/j.acalib.2024.102899>

Case, D. O., & Given, L. M. (2016). *Looking for information: A survey of research on information seeking, needs, and behavior*. Emerald.

Celik, I., Kontkanen, S., Laru, J., & Dalyanci, A. A. (2025). Co-constructing adaptive lesson plans with GenAI: Pre-service teachers' Intelligent-TPACK and prompt engineering strategies. *Computers & Education*. , Article 105485. <https://doi.org/10.1016/j.compedu.2025.105485>

Cengiz, S., & Peker, A. (2025). Generative artificial intelligence acceptance and artificial intelligence anxiety among university students: The sequential mediating role of attitudes toward artificial intelligence and literacy. *Current Psychology*, 44(9), 7991–8000. <https://doi.org/10.1007/s12144-025-07433-7>

Cosgrove, J., & Cachia, R. (2025). *DigComp 3.0: European digital competence framework*. European Commission.

Criado, J. I., Alcaide-Muñoz, L., & Liarte, I. (2025). Two decades of public sector innovation: Building an analytical framework from a systematic literature review of types, strategies, conditions, and results. *Public Management Review*, 27(3), 623–652. <https://doi.org/10.1080/14719037.2023.2254310>

Dauer, J. M., Sorensen, A. E., & Jimenez, P. C. (2022). Using structured decision-making in the classroom to promote information literacy in the context of decision-making. *Journal of College Science Teaching*, 51(6), 75–82. <https://doi.org/10.1080/0047231X.2022.12315652>

De Meulemeester, A., De Maeseneer, J., De Maeyer, S., Peleman, R., & Buysse, H. (2018). Information literacy self-efficacy of medical students: A longitudinal study. In *European conference on information literacy* (pp. 264–272). Cham: Springer International Publishing.

Dinsmore, D. L., & Fryer, L. K. (2026). What does current genAI actually mean for student learning? *Learning and Individual Differences*, 125, Article 102834. <https://doi.org/10.1016/j.lindif.2025.102834>

Draxler, F., Werner, A., Lehmann, F., Hoppe, M., Schmidt, A., Buschek, D., & Welsch, R. (2024). The AI ghostwriter effect: When users do not perceive ownership of AI-generated text but self-declare as authors. *ACM Transactions on Computer-Human Interaction*, 31(2), 1–40. <https://doi.org/10.1145/3637875>

Elmborg, J. (2006). Critical information literacy: Implications for instructional practice. *The Journal of Academic Librarianship*, 32(2), 192–199. <https://doi.org/10.1016/j.acalib.2005.12.004>

Elyoseph, Z., Hadar Shoval, D., & Levkovich, I. (2024). Beyond personhood: Ethical paradigms in the generative artificial intelligence era. *The American Journal of Bioethics*, 24(1), 57–59. <https://doi.org/10.1080/15265161.2023.2278546>

Elyoseph, Z., & Levkovich, I. (2024). Comparing the perspectives of GenAI, mental health experts, and the general public on schizophrenia recovery: Case vignette study. *JMIR Mental Health*, 11(1), Article e53043. <https://doi.org/10.2196/53043>

Figenschou, T., Li-Ying, J., Tanner, A., & Bogers, M. (2024). Open innovation in the public sector: A literature review on actors and boundaries. *Technovation*, 131, Article 102940. <https://doi.org/10.1016/j.technovation.2023.102940>

Geroimenko, V. (2026). Feedback and iterative refinement: Improving AI responses through interaction cycles. In *Beyond and after prompt engineering: The future of AI communication*. Cham: Springer. https://doi.org/10.1007/978-3-032-04569-0_11

Gkinko, L., & Elbanna, A. (2022). Hope, tolerance and empathy: Employees' emotions when using an AI-enabled chatbot in a digitalised workplace. *Information Technology & People*, 35(6), 1714–1743. <https://doi.org/10.1108/ITP-04-2021-0328>

Hadar-Shoval, D., Asraf, K., Shinan-Altman, S., Elyoseph, Z., & Levkovich, I. (2024). Embedded values-like shape ethical reasoning of large language models on primary care ethical dilemmas. *Heliyon*, 10(18), Article e38056. <https://doi.org/10.1016/j.heliyon.2024.e38056>

Haesevoets, T., Verschuere, B., & Roets, A. (2025). AI adoption in public administration: Perspectives of public sector managers and public sector non-managerial employees. *Government Information Quarterly*, 42(2). <https://doi.org/10.1016/j.giq.2025.102029>

Hong, S. J. (2025). What drives AI-based risk information-seeking intent? Insufficiency of risk information versus (un)certainly of AI chatbots. *Computers in Human Behavior*, 162, Article 108460. <https://doi.org/10.1016/j.chb.2024.108460>

Hornberger, M., Bewersdorff, A., Schiff, D. S., & Nerdel, C. (2025). Development and validation of a short AI literacy test (AILIT-S) for university students. *Computers in Human Behavior: Artificial Humans*, 5, Article 100176. <https://doi.org/10.1016/j.chbah.2025.100176>

Hou, C., Zhu, G., Sudarshan, V., Lim, F. S., & Ong, Y. S. (2025). Measuring undergraduate students' reliance on Generative AI during problem-solving: Scale development and validation. *Computers & Education*, 105329. <https://doi.org/10.1016/j.compedu.2025.105329>

Hutt, S., DePiro, A., Wang, J., Rhodes, S., Baker, R. S., Hieb, G., & Mills, C. (2024). Feedback on feedback: Comparing classic natural language processing and GenAI to evaluate peer feedback. In *Proceedings of the 14th learning analytics and knowledge conference* (pp. 55–65). <https://doi.org/10.1145/3636555.3636850>

Huvila, I., & Gorichanaz, T. (2025). Trends in information behavior research, 2016–2022: An annual review of information science and technology (ARIST) paper. *Journal of the Association for Information Science and Technology*, 76(1), 216–237. <https://doi.org/10.1002/asi.24943>

John, K., & Tater, B. (2025). Reframing the information literacy framework to identify misinformation and disinformation. *The Serials Librarian*, 86(1-2), 29–55. <https://doi.org/10.1080/0361526X.2025.2459765>

- Johnson, J. G., Peralta, M., Kaur, M., Huang, R. S., Zhao, S., Guan, R., & Nebeling, M. (2025). Exploring collaborative GenAI agents in synchronous group settings: Eliciting team perceptions and design considerations for the future of work. *Proceedings of the ACM on Human-Computer Interaction*, 9(7), 1–33. <https://doi.org/10.1145/3757595>
- Jones, W., Capra, R., Diekema, A., Teevan, J., Pérez-Quinones, M., Dinneen, J. D., & Hemminger, B. (2015). "For telling" the present: Using the Delphi method to understand personal information management practices. In *Proceedings of the 33rd annual ACM conference on human factors in computing systems* (pp. 3513–3522). <https://doi.org/10.1145/2702123.2702523>
- Jorolan, J., Cabillo, F., Batucan, R. R., Camansi, C. M., Etoquilla, A., Gapo, J., & Gonzales, G. (2025). Development and validation of moral ascendancy and dependency in AI integration scale (MAD-AI) for teachers. *Computers & Education*, 235, Article 105346. <https://doi.org/10.1016/j.compedu.2025.105346>
- Keshavarz, H., Esmailie Givi, M. R., Vafaeian, A., & Khademan, M. (2017). Information literacy self-efficacy dimensions of post graduate students: Validating a Persian version scale. *Libri - International Journal of Libraries and Information Services*, 67(1), 75–86. <https://doi.org/10.1515/libri-2016-0092>
- Koch, M. J., Wienrich, C., Straka, S., Latoschik, M. E., & Carolus, A. (2024). Overview and confirmatory and exploratory factor analysis of AI literacy scale. *Computers and Education: Artificial Intelligence*, 7, Article 100310. <https://doi.org/10.1016/j.caeai.2024.100310>
- Kuhlthau, C. C. (1991). Inside the search process: Information seeking from the user's perspective. *Journal of the American Society for Information Science*, 42(5), 361–371. [https://doi.org/10.1002/\(SICI\)1097-4571\(199106\)42:5<361::AID-ASIS6>3.0.CO;2-%23](https://doi.org/10.1002/(SICI)1097-4571(199106)42:5<361::AID-ASIS6>3.0.CO;2-%23)
- Kuhlthau, C. C., Heinström, J., & Todd, R. J. (2008). The 'information search process' revisited: Is the model still useful. *Information Research*, 13(4), 13–14.
- Lee, H. Y., Chen, P. H., Lin, C. J., Huang, Y. M., & Wu, T. T. (2025). Leveraging ChatGPT for personalized reflective learning in programming education: Effects on self-efficacy, higher-order thinking, and project implementation skills. *Education and Information Technologies*, 30(17), 24815–24854. <https://doi.org/10.1007/s10639-025-13733-z>
- Levkovich, I. (2025). Harnessing large language models for identification and treatment of obsessive-compulsive disorder. *Computers in Human Behavior: Artificial Humans*. <https://doi.org/10.1016/j.chbah.2025.100212>
- Levkovich, I., Shinar-Altman, S., & Elyoseph, Z. (2024). Can large language models be sensitive to culture suicide risk assessment? *Journal of Cultural Cognitive Science*, 8(3), 275–287. <https://doi.org/10.1007/s41809-024-00151-9>
- Li, R., Ouyang, S., & Lin, J. (2025). Mediating effect of AI attitudes and AI literacy on the relationship between career self-efficacy and job-seeking anxiety. *BMC Psychology*, 13(1), 454. <https://doi.org/10.1186/s40359-025-02757-2>
- Li, J., & Yuan, Q. (2026). The impact of AI-generated knowledge on user knowledge avoidance in OKCs: The moderating role of AI literacy and platform trust. *Journal of Knowledge Management*, 1–25. <https://doi.org/10.1108/JKM-02-2025-0231>
- Lintner, T. (2024). A systematic review of AI literacy scales. *Npj Science of Learning*, 9(1), 50. <https://doi.org/10.1038/s41539-024-00264-4>
- Liu, Y., & Du, Y. (2025). The effect of generative AI ethics on users' continuous usage intentions: A PLS-SEM and fsQCA approach. *International Journal of Human-Computer Interaction*, 1–12. <https://doi.org/10.1080/10447318.2025.2465861>
- Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI conference on human factors in computing systems* (pp. 1–16). <https://doi.org/10.1145/3313831.3376727>
- Lv, D., Sun, R., Zhu, Q., & Qin, S. (2025). Preconceived beliefs, different reactions: Alleviating user switching intentions in service failures through priming GenAI beliefs. *BMC Psychology*, 13(1), 552. <https://doi.org/10.3390/jintelligence13070078>
- Maleki, N., Padmanabhan, B., & Dutta, K. (2024). AI hallucinations: A misnomer worth clarifying. In *2024 IEEE conference on artificial intelligence (CAI)* (pp. 133–138). IEEE. <https://doi.org/10.1109/CAI59869.2024.00033>
- Messeri, L., & Crockett, M. J. (2024). Artificial intelligence and illusions of understanding in scientific research. *Nature*, 627(8002), 49–58. <https://doi.org/10.1038/s41586-024-07146-0>
- Naveed, M. A., & Mahmood, M. (2022). Correlatives of business students' perceived information literacy self-efficacy in the digital information environment. *Journal of Librarianship and Information Science*, 54(2), 294–305. <https://doi.org/10.1177/09610006211014277>
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, Article 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
- OECD. (2025). Empowering learners for the age of AI: An AI literacy framework for primary and secondary education. Retrieved January 2026 from <https://ailiteracyframework.org/>.
- Ooi, K. B., Tan, G. W. H., Al-Emran, M., Al-Sharafi, M. A., Capatina, A., Chakraborty, A., & Wong, L. W. (2025). The potential of generative artificial intelligence across disciplines: Perspectives and future directions. *Journal of Computer Information Systems*, 65(1), 76–107. <https://doi.org/10.1080/08874417.2023.2261010>
- Radanliev, P. (2025). AI ethics: Integrating transparency, fairness, and privacy in AI development. *Applied Artificial Intelligence*, 39(1), Article 2463722. <https://doi.org/10.1080/08839514.2025.2463722>
- Ravuri, B., & Mardis, M. A. (2025). AI on the shoulders of giants: Using kuhlthau's information search process to improve AI support for information-seeking. *Library Trends*, 73(3), 267–296. <https://doi.org/10.1353/lib.2025.a961195>
- Rezaei, M., Pironti, M., & Quaglia, R. (2025). AI in knowledge sharing, which ethical challenges are raised in decision-making processes for organisations? *Management Decision*, 63(10), 3369–3388. <https://doi.org/10.1108/MD-10-2023-2023>
- Savolainen, R. (2008). Source preferences in the context of seeking problem-specific information. *Information Processing & Management*, 44(1), 274–293. <https://doi.org/10.1016/j.ipm.2007.02.008>
- Schuster, A. M., & Cotten, S. R. (2024). Differences between employed and retired older adults in information and communication technology use and attitudes. *Work, Aging and Retirement*, 10(1), 38–45. <https://doi.org/10.1093/workar/waac025>
- Serap Kurbanoglu, S., Akkoyunlu, B., & Umay, A. (2006). Developing the information literacy self-efficacy scale. *Journal of Documentation*, 62(6), 730–743. <https://doi.org/10.1108/00220410610714949>
- Shen, L., Li, H., Wang, Y., Xie, X., & Qu, H. (2025). Prompting generative AI with interaction-augmented instructions. In *Proceedings of the extended abstracts of the CHI conference on human factors in computing systems* (pp. 1–9). <https://doi.org/10.1145/3706599.3720080>
- Shin, D., Koerber, A., & Lim, J. S. (2025). Impact of misinformation from generative AI on user information processing: How people understand misinformation from generative AI. *New Media & Society*, 27(7), 4017–4047. <https://doi.org/10.1177/14614448241234040>
- Sigot, M., & Tassoti, S. (2025). An investigation of change in prompting strategies in a semester-long course on the use of GenAI. *Journal of Chemical Education*, 102(6), 2507–2513. <https://doi.org/10.1021/acs.jchemed.4c01287>
- Singhal, A., Neveditsin, N., Tanveer, H., & Mago, V. (2024). Toward fairness, accountability, transparency, and ethics in AI for social media and health care: Scoping review. *JMIR Medical Informatics*, 12(1), Article e50048. <https://doi.org/10.2196/50048>
- Sun, Y., Jang, E., Ma, F., & Wang, T. (2024). GenAI in the wild: Prospects, challenges, and strategies. In *Proceedings of the 2024 CHI conference on human factors in computing systems* (pp. 1–16). <https://doi.org/10.1145/3613904.3642160>
- Trixa, J., & Kaspar, K. (2024). Information literacy in the digital age: Information sources, evaluation strategies, and perceived teaching competences of pre-service teachers. *Frontiers in Psychology*, 15, Article 1336436. <https://doi.org/10.3389/fpsyg.2024.1336436>
- Ulfert-Blank, A. S., & Schmidt, I. (2022). Assessing digital self-efficacy: Review and scale development. *Computers & Education*, 191, Article 104626. <https://doi.org/10.1016/j.compedu.2022.104626>
- Vuorikari, R., Kluzer, S., & Punie, Y. (2022). DigComp 2.2: The digital competence framework for citizens-with new examples of knowledge, skills and attitudes. Retrieved December 2025: <https://vibe.cityoflearning.eu/storage/content/u3h4e3jb5t4lpvm44vs5vmrn3om3kgq1.pdf>.
- Vuorikari, R., Pokropek, A., & Muñoz, J. C. (2025). Enhancing digital skills assessment: Introducing compact tools for measuring digital competence. *Technology, Knowledge and Learning*, 1–28. <https://doi.org/10.1007/s10758-025-09825-x>
- Wei, X., Kumar, N., & Zhang, H. (2025). Addressing bias in GenAI: Challenges and research opportunities in information management. *Information & Management*, 62(2), Article 104103. <https://doi.org/10.1016/j.im.2025.104103>
- Williamson, S. M., & Prybutok, V. (2024). The era of artificial intelligence deception: Unraveling the complexities of false realities and emerging threats of misinformation. *Information*, 15(6), 299. <https://doi.org/10.3390/info15060299>
- Wilson, T. D. (1999). Models in information behaviour research. *Journal of Documentation*, 55(3), 249–270. <https://doi.org/10.1108/EUM0000000007145>
- Zhang, H., & Wang, Q. (2025). How could GenAI work on in-service teachers' knowledge building process? An empirical study based on epistemic network analysis. *International Journal of Educational Technology in Higher Education*, 22(1), 47. <https://doi.org/10.1186/s41239-025-00544-y>
- Zhu, W., Huang, L., Zhou, X., Li, X., Shi, G., Ying, J., & Wang, C. (2025). Could AI ethical anxiety, perceived ethical risks and ethical awareness about AI influence university students' use of generative AI products? An ethical perspective. *International Journal of Human-Computer Interaction*, 41(1), 742–764. <https://doi.org/10.1080/10447318.2024.2323277>