



www.ijte.net

Ethical and Social Risk Awareness in Generative AI (GenAI): The Role of Mindset and GenAI Literacy

Sahin Gokcearslan 
Hacettepe University, Türkiye

Hatice Yildiz Durak 
Necmettin Erbakan University, Türkiye

Mustafa Serkan Gunbatar 
Van Yuzuncuyıl University, Türkiye

Nilufer Atman Uslu 
Dokuz Eylul University, Türkiye

Aysun Nuket Elci 
Dokuz Eylul University, Türkiye

To cite this article:

Gokcearslan, S., Yildiz Durak, H., Gunbatar, M.S., Atman Uslu, N., & Elci, A.N. (2025). Ethical and social risk awareness in generative AI (GenAI): The role of mindset and GenAI literacy. *International Journal of Technology in Education (IJTE)*, 8(4), 914-937. <https://doi.org/10.46328/ijte.1565>

The International Journal of Technology in Education (IJTE) is a peer-reviewed scholarly online journal. This article may be used for research, teaching, and private study purposes. Authors alone are responsible for the contents of their articles. The journal owns the copyright of the articles. The publisher shall not be liable for any loss, actions, claims, proceedings, demand, or costs or damages whatsoever or howsoever caused arising directly or indirectly in connection with or arising out of the use of the research material. All authors are requested to disclose any actual or potential conflict of interest including any financial, personal or other relationships with other people or organizations regarding the submitted work.



This work is licensed under a Creative Commons Attribution-NonCommercial-ShareAlike 4.0 International License.

Ethical and Social Risk Awareness in Generative AI (GenAI): The Role of Mindset and GenAI Literacy

Sahin Gokcearslan, Hatice Yildiz Durak, Mustafa Serkan Gunbatar, Nilufer Atman Uslu, Aysun Nuket Elci

Article Info

Article History

Received:

24 February 2025

Accepted:

27 June 2025

Keywords

Generative artificial
intelligence

Ethical and social risks harm
awareness

GenAI literacy

Mindset

Measurement tool

Abstract

GenAI's advanced natural language processing capabilities will revolutionize numerous areas, ranging from a paradigm shift in education to the economy. Along with the positive aspects of GenAI, ethical and social risks are also one of the negative aspects that attract attention in the literature. The purpose of this study is to test the role of mindset and GenAI literacy in university students' awareness of ethical and social risk with structural equation modeling. According to the research results, the developed awareness of ethical and social risk of generative artificial intelligence (GenAI) and the adapted GenAI literacy scale are valid and reliable. GenAI literacy variable does not have a significant effect on ethical and social risk awareness towards GenAI. Mindset should be taken into account in activities for both variables.

Introduction

Generative AI (GenAI) brings many advantages in terms of personalized, fast, and effective content creation today. GenAI's improved capabilities are projected to potentially automate up to two-thirds of work activities by 2023, with a forecast of half of today's work activities being automated between 2030 and 2060, primarily attributed to GenAI's enhanced natural language abilities, impacting higher-skilled and higher-paying occupations, and potentially increasing labor productivity by 0.1 to 0.6 percent annually over the next decade to two decades (Chui et al, 2023). The introduction of state-of-the-art GenAI is poised to transform society by revolutionizing how we live, work, learn, and communicate, as it can create diverse multimodal content and facilitate problem-solving, while also offering assistance to students and teachers in various educational tasks (Fui-Hoon Nah et al. 2023).

GenAI uses a language model, and this model has many ethical and social risks and harms as well as positive aspects. Weidinger et al. (2021) examines these risks under six classifications. NLP can perpetuate discrimination, exclusion, and toxicity by reinforcing social stereotypes and unfair discrimination, potentially leading to the exclusion of certain languages and social groups, exacerbating inequalities and manifesting in various forms of harm such as information hazards, misinformation, malicious uses, human-computer interaction issues, and

environmental and access concerns related to automation (Weidinger et al., 2021). Wach et al. (2023) identified Gen AI's risk as including the urgent need for regulation, poor quality and lack of control, job losses, personal data violation, social manipulation, socio-economic inequalities, and AI technostress. Individuals need to be aware of these negative consequences and ethical issues when using GenAI. However, it is noteworthy that the studies on Ethical and social risk harm awareness are limited. GenAI literacy provides an important starting point in the formation of this awareness. AI literacy refers to a new set of technological competencies required for people to use AI ethically and effectively in their daily lives by combining social and technical skills (Pinski & Benlian, 2023; Ng et al., 2022). The concept of AI literacy can be addressed in the context of fields such as Education, Information & Knowledge Processing, Human Resources & Industrial Relations (Çelebi et al., 2023). Therefore, different dimensions may come to the fore in different disciplines. However, almost all researchers emphasize the concept of ethics when defining AI literacy. For example, while defining AI literacy, Wang et al. (2023) took into account the dimensions of *awareness, use, evaluation, and ethics*; Ng et al. (2021) *know and understand AI, use and apply AI, evaluate and create AI, and AI ethics*; Kong & Zhang (2021) *knowing concepts of AI, daily use of AI, ethical use of AI*. One of the reasons for such a clear emphasis on the concept of ethics is the ethical and social risks that may arise when using this technology. Therefore, individuals who are AI literate have an understanding of the ethical and responsible use of this technology.

Although there is no direct research on mindset and AI ethics, it is found in the literature that having a growth mindset is an important quality for individuals to be ethical learners (Chugh & Kern, 2016) and it can be argued that examining the role of individuals' mindsets in ethical and social risk harm awareness may yield meaningful findings in deepening the understanding on this subject. According to the mindset theory, there are two types of mindset: (1) Growth mindset (2) Fixed mindset (Dweck, 2006). Individuals with a fixed mindset, intelligence is static. They tend to avoid challenges, give up easily in the face of obstacles, do not show much effort, ignore criticism, and perceive the success of others as a threat to themselves. Despite this, according to individuals with a growth mindset, intelligence can be developed. They tend to embrace challenges, persist in the face of obstacles, see effort as the part of mastery, learn from criticism, learn from others' success and be inspired (The Open University, 2023). Mindset variable emerges as an important determinant for individuals' learning and development (Yılmaz, 2022) and has reflections on the communication processes with GenAI (Guo, Zhong & Chu, 2023). Mindset shapes people's responses in their interactions with AI-powered robots. Individuals with a growth mindset tend to react more positively than individuals with a fixed mindset. This is because they are less concerned about revealing their limitations of mind and being surpassed by robots (Dang & Liu, 2022). Individuals with a Growth mindset have an advantage in adapting to this new technology (Farrow, 2021). As a matter of fact, there are several studies showing the role of individuals' mindsets in being at risk in their technology use behaviors (Lee-Won et al., 2020) or having positive perceptions about new technologies (Dang & Liu, 2022).

It can be claimed that there is a need for studies examining the effect of mindset on the ethical and social risk harm awareness of Gen AI. Based on these points, the current study aims to examine the effect of fixed and growth mindset on Gen AI literacy's awareness of possible ethical and social risks and harms, and makes an attempt to contribute to the development of nomological networks on this subject.

Literature Review

Generative AI (GenAI) and Ethical and Social Risk Harm Awareness

Generative AI (GenAI) systems, such as ChatGPT, generate human-like responses using large datasets and statistical models (Mohamed, 2023; Tirado-Olivares et al., 2023). These systems rely on probabilistic associations rather than factual understanding, making it difficult to distinguish truth from falsehood (Harrer, 2023; Sison et al., 2023). As such, they are often described as “stochastic parrots” (Harrer, 2023; Sison et al., 2023). The outputs may contain bias, reinforce power structures, or mislead users (Weinberg, 2022).

The ethical risks of GenAI can be considered in two dimensions: (1) how GenAI generates information, and (2) how people use this information. International organizations such as WHO (World Health Organization) have identified major risks, including accountability, fairness, privacy, transparency, explainability, and value alignment (Harrer, 2023). Misuses include academic dishonesty, misinformation, and criminal facilitation (Sison et al., 2023). To guide ethical use, frameworks like PAPA-Privacy, Accuracy, Property, Accessibility—are useful (Mason, 1986; Niederman & Baker, 2023). UNESCO's global guidelines also stress the importance of ethics in educational use (Miao & Holmes, 2023).

GenAI is now widely used by students (Tlili et al., 2023), patients (Nov et al., 2023), researchers (Lin, 2023), and developers (Denny et al., 2023). While it can enhance learning (Bahroun et al., 2023), maximizing its benefits requires clear ethical guidelines (Májovský et al., 2023; Akkaş et al., 2024) and public awareness (Su et al., 2023). Policies must address pedagogical, governance, and operational aspects (Chan, 2023; Lim et al., 2023), supported by ongoing research (Panthier & Gatinel, 2023).

Cheng and Liu (2023) highlight key ethical principles: accountability, privacy, transparency, fairness, security, safety, non-discrimination, accessibility, explainability, and responsibility. In education, Lim et al. (2023) outline four paradoxes: GenAI is a ‘friend’ yet a ‘foe’, ‘capable’ yet ‘dependent’, ‘accessible’ yet ‘restrictive’, and ‘popular’ even when ‘banned’. These tensions show the need for balanced, ethical integration rather than outright restriction.

This study draws on Weidinger et al. (2021), identifying six categories of harm: (1) discrimination and toxicity, (2) information hazards, (3) misinformation, (4) malicious use, (5) human-computer interaction issues, and (6) automation, access, and environmental harms. These categories inform the measurement of risk awareness in this research. In this context, GenAI literacy—defined as the ability to critically understand, evaluate, and use generative AI systems—can serve as a protective factor, enhancing users’ sensitivity to these risks. This increased awareness fosters more responsible and informed decision-making, thereby positively influencing their ability to identify and mitigate potential harms. Similarly, Rozak and Karman (2025) emphasize that digital literacy is critical in enabling users to recognize harmful content and ethical issues in the online environment, be aware of the risks associated with them, and make responsible and ethical usage decisions. Abuadas and Albikawi (2025) find that high digital and AI literacy increases the likelihood that users are better able to recognize and notice such ethical and social harms, and that this awareness plays an important role in responsible decision-making and risk

mitigation. Strauß (2021) emphasizes that technological systems need to be critically evaluated in a social context to better understand the potential harms of AI systems and states that such awareness is only possible with high digital and AI literacy. Therefore, it is posited that individuals with higher GenAI literacy are more likely to recognize and be aware of the ethical and social harms associated with its use.

H1: GenAI literacy positively influences ethical and social risk awareness.

According to Dweck (2006), individuals adopt either a fixed or growth mindset. Those with a fixed mindset believe intelligence is static, avoid challenges, give up easily, ignore feedback, and perceive others' success as a threat. In contrast, individuals with a growth mindset believe intelligence can develop, embrace challenges, persist despite setbacks, and learn from feedback and others' success (The Open University, 2023). This mindset influences how people handle negative experiences—while fixed-minded individuals dwell on failure, those with a growth mindset focus on improvement (Dweck, 2006). Individuals who believe their abilities will be improved through effort and action will likely take action (Kaltenegger, 2024). These characteristics suggest that a fixed mindset limits individuals' motivation and openness toward learning and growth, which are essential in acquiring GenAI literacy.

H2a: Fixed Mindset has a negative influence on GenAI Literacy.

Mindset also plays a role in interactions with AI. People with a growth mindset respond more positively to AI technologies because they are less threatened by the possibility of being outperformed (Dang & Liu, 2022; Kaltenegger, 2024). In future scenarios involving AI-related layoffs, growth-minded individuals demonstrated stronger adaptability and future literacy, engaging in mutual learning and problem-solving, whereas fixed-minded individuals expressed fear and sadness (Farrow, 2021). Such emotional responses and reduced engagement may limit fixed-minded individuals' awareness of ethical and social risks related to GenAI. A fixed mindset, which views abilities as innate and unchangeable, can hinder reflection on the broader societal impacts of emerging technologies.

H2b: Fixed Mindset has a negative influence on GenAI Risk Awareness for Ethical and Social Harms.

Similarly, in learning environments, fixed-minded individuals focus only on whether an answer is right, while growth-minded ones explore the reasoning to deepen their understanding (Dweck, 2006). Given these distinctions, mindset significantly affects individual learning and communication across contexts, including education, professional life, and interpersonal relationships (Dweck, 2006; Yilmaz, 2022). As GenAI tools like ChatGPT have become widely accessible (Guo, Zhong & Chu, 2023), people with a growth mindset may be more adaptable and capable of using GenAI ethically (Mehan, 2023; Farrow, 2021). Farrow (2020) reports that adapting to AI-related innovations is a necessary component that drives AI literacy. These findings suggest that individuals with a growth mindset are more capable of acquiring GenAI literacy because they are willing to explore, experiment, and reflect.

H3a: Growth Mindset has a positive influence on GenAI Literacy.

Although direct studies on mindset and GenAI ethics are scarce, the literature highlights a growth mindset as crucial for ethical learning (Chugh & Kern, 2016) and for success in digital work environments influenced by AI

(Athota, 2021). Since mindset has not been sufficiently addressed in studies on GenAI literacy and ethical and social risks related to GenAI, hypothesis testing can lay the foundation for future studies. Expert reflections confirm that a growth mindset supports ethical decision-making (Norman, Mayowski & Fine, 2021). Chiu (2024) noted that students' passive reliance on GenAI tools without critical engagement may hinder their understanding of how these tools function, which aligns with patterns typically associated with a fixed mindset. The development of AI literacy should aim to enhance individuals' ethical perspectives and change their mindsets (Li & Kim, 2024). Therefore, growth-minded individuals are expected to gain more from GenAI interaction, especially in terms of ethical awareness (Su, Lin & Lai, 2023).

H3b: Growth Mindset has a positive influence on GenAI Risk Awareness for Ethical and Social Harms.

GenAI Literacy, GenAI Ethical and Social Risk Harm Awareness and Mindset

There is a need for more practical and problem-oriented analytical perspectives on the risks of artificial intelligence to overcome the focus on impractical ethical issues and technocratic approaches (Strauß, 2021). Accordingly, it can be claimed that individuals' AI literacy is an important starting point in increasing their awareness of potential risks. According to Kong et al. (2023), AI literacy has cognitive, affective and socio-cultural dimensions. In this context, the sociocultural dimension of AI literacy includes awareness of ethical issues (Kong et al., 2023). Ng et al. (2021) presented a framework that includes four aspects of AI literacy (know and understand, use and apply, evaluate and create, and ethical issues). Also, Kong et al. (2013) found that a programme for AI literacy has positively affected undergraduates' ethical awareness. From these points, the following hypothesis is formulated:

H4: GenAI Literacy has a positive influence on GenAI Risk Awareness for Ethical and Social Harms.

Method

Participants

The study was conducted in two phases. In Study 1, 66.7% of the 297 participants were female and 33.3% were male. Their mean age was 21.75 years. 85.9% were associate and undergraduate students, 14.1% were graduate and formation students. Of the 441 participants who participated in the study in phase 2, 65.3% were female and 34.7% were male. 68.3% of the participants were studying at the faculties of educational sciences, 17.7% at the faculties of sport sciences and 14.1% at the faculties of basic sciences and engineering. 85.3% of the students were undergraduate and 14.7% were graduate students. In the selection of the participants, convenience sampling method was used. According to Kılıç (2013), this sampling method is one of the non-probability sampling techniques and enables the researcher to collect data from individuals or groups that are easily accessible and practical to reach. In this method, which is preferred for practical reasons such as speeding up the research, reducing costs and saving time, the selection of the individuals forming the sample is not random, but based on accessibility. In this context, in order to facilitate the data collection process, the researcher included individuals who are close to their own environment, easy to reach and who can participate voluntarily. Convenience sampling is a frequently used method, especially in preliminary research, pilot studies or in cases of limited resources.

Measurement Tools

GenAI Ethical and Social Harm Risk Awareness Scale

The researchers developed the risk awareness scale for ethical and social harms of productive artificial intelligence from the construct that Weidinger et al. (2021) addressed in 6 dimensions and supported these dimensions with references (see Appendix A and B for Turkish and English versions). The feature in each dimension related to the construct was supported by Weidinger et al (2021) with references as a result of a comprehensive literature review. For example, one of the features of the Human computer Interaction Harms dimension, "Anthropomorphizing systems can lead to overreliance or unsafe use" Kim and Sundar (2012), and "Create avenues for exploiting user trust to obtain private information" (Ischen et al. (2019; Lewis et al., 2017; Van den Broeck et al., 2019).

The item pool was developed by the current study's researchers. The developed items were presented to 1 language expert, 1 measurement expert, 5 educational technology and 2 informatics field experts. The items were finalized after expert opinion. No items were removed from the item pool.

Generative AI Literacy Scale

The measurement tool developed by Wang et al. (2023) to measure AI literacy was adapted to generative AI. An item pool of 65 items was created for the construct consisting of awareness, usage, evaluation and ethics dimensions. 31 items were subjected to reliability and validity tests on two separate samples. A 12-item instrument was obtained to measure AI literacy. The instrument was found to be significantly related to digital literacy, attitudes towards robots, and users' daily use of AI. A high score on the instrument indicates a high level of AI literacy. 12 items provide a general measure of AI literacy. These general statements were subjected to confirmatory factor analysis with GenAI adaptation.

Mindset Scale

Exploratory and confirmatory factor analysis was applied with the participation of 1145 students to measure the Mindset type of students aged 14-22. The 26-item measurement tool was reduced to 19 items (four sub-dimensions) as a result of exploratory factor analysis. The scale is in 5-point Likert format. In line with the literature, these dimensions were named as Procrastination, Stability of Belief, Belief in Development and Effort. The four-factor structure of the scale was confirmed by confirmatory factor analysis. In addition, it was found that the differences between the means of the lower and upper groups, which constituted 27% of the scale items, were significant. The internal consistency coefficient was found to be 0.723 for the Fixed Mindset dimension and 0.714 for the Growth Mindset dimension (Yılmaz, 2022).

Research Ethics

Before the study, participants were informed about its focus and scope and the steps to be followed to ensure data confidentiality. Participants' information was anonymized, and the dataset was stored on a computer to which only

the research team had access. Participation in the study was entirely voluntary. Participants were given voluntary consent forms and informed that they could withdraw from the study at any time.

Results

Descriptive Findings

The mean, standard deviation, skewness and kurtosis statistics and standard error values of the items in this scale are presented in Table 1.

Table 1. Mean, Standard Deviation, Kurtosis and Skewness

Items	Mean	Standard Deviation	Skewness		Kurtosis	
			Statistic	Std. Error	Statistic	Std. Error
A1	5.38	1.208	-0.202	0.141	-0.981	0.282
A2	3.42	1.733	0.411	0.141	-0.768	0.282
A3	4.39	1.321	0.233	0.141	-0.461	0.282
A4	4.77	1.333	-0.063	0.141	-0.588	0.282
A5	3.67	1.592	0.3	0.141	-0.647	0.282
A6	5.46	1.205	-0.325	0.141	-0.844	0.282
A7	5.25	1.226	-0.235	0.141	-0.844	0.282
A8	5.32	1.163	-0.185	0.141	-0.778	0.282
A9	5.05	1.155	0.04	0.141	-0.73	0.282
A10	4.96	1.538	-0.392	0.141	-0.603	0.282
A11	2.4	1.528	1.064	0.141	0.384	0.282
A12	4.96	1.505	-0.176	0.141	-1.037	0.282
A13	2.55	1.127	0.122	0.141	-0.839	0.282
A14	2.37	1.016	0.045	0.141	-1.135	0.282
A15	2.05	0.982	0.789	0.141	0.311	0.282
A16	2.57	1.11	0.144	0.141	-0.795	0.282
A17	2.69	0.944	0.323	0.141	-0.033	0.282
A18	3.15	1.096	-0.064	0.141	-0.665	0.282
A19	2.73	0.973	0.133	0.141	-0.291	0.282
A20	2.87	1.045	0.223	0.141	-0.397	0.282
A21	2.77	0.979	0.053	0.141	-0.614	0.282
A22	2.74	1.054	0.012	0.141	-0.611	0.282
A23	2.76	1.08	0.222	0.141	-0.573	0.282
A24	2.51	1.027	0.467	0.141	-0.185	0.282
A25	2.28	1.007	0.448	0.141	-0.371	0.282
A26	2.71	1.116	0.114	0.141	-0.686	0.282
A27	2.59	1.112	0.368	0.141	-0.51	0.282

Items	Mean	Standard Deviation	Skewness		Kurtosis	
			Statistic	Std. Error	Statistic	Std. Error
A28	2.66	1.001	0.174	0.141	-0.439	0.282
A29	2.56	1.042	0.267	0.141	-0.403	0.282
A30	2.49	1.007	0.302	0.141	-0.337	0.282
A31	2.43	1.011	0.406	0.141	-0.275	0.282
A32	2.56	1.108	0.398	0.141	-0.442	0.282

According to Table 1, the mean scores of the items are between 2.05 and 5.46. The standard deviations of the items vary between 1.773 and 0.944. The skewness and kurtosis values calculated for each item are between +1.5 and -1.5. When the normal distribution curve of the items was analyzed, it was concluded that the scores were normally distributed.

Confirmatory Factor Analysis

The theoretical structure of the scale was developed by Weidinger et al (2021) through a comprehensive literature review, and therefore, confirmatory factor analysis was performed first. The fit index values of the confirmatory factor analysis were found as [χ^2 /sd=2.93, RMSEA=.08, GFI=.90, CFI=.96, NFI=.93, NNFI=.95]. It showed acceptable and excellent fit values for the fit indices examined in order to reveal the adequacy of the model. This reveals that the fit level of the six-factor model obtained from CFA is adequate. These values and the ranges accepted in the literature are presented in Table 2.

Table 2. CFA Fit Indices Values

Goodness of fit criteria	Recommended values	Acceptable fit	Reference	Observed values	Fit situations
χ^2 /df	$0 \leq \chi^2 /sd \leq 2$	$2 \leq \chi^2 /sd \leq 3$	(Çelik &	2.93	Acceptable
RMSEA	$.00 \leq RMSEA \leq .05$	$.05 \leq RMSEA \leq .08$	Yılmaz, 2013;	0.08	Acceptable
GFI	$.95 \leq GFI \leq 1.00$	$.90 \leq GFI \leq .95$	Çokluk et al.,	0.90	Acceptable
CFI	$.95 \leq CFI \leq 1.00$	$.90 \leq CFI \leq .95$	2012; Schumacker	0.96	Good
NFI	$.95 \leq NFI \leq 1.00$	$.90 \leq NFI \leq .95$	& Lomax, 2004)	0.93	Acceptable
NNFI	$.95 \leq NNFI \leq 1.00$	$.90 \leq NNFI \leq .95$		0.95	Good

Confirmatory factor analysis resulted in standardized factor loadings and item structures are presented in Figure 1. According to the figure, factor loadings ranged between .44 and .93 for all items. According to the t-test findings, all connections are statistically significant. Two field experts were consulted to evaluate whether the items with factor loadings below 0.7 should be removed from the scale. The field experts stated that the content of the scale was important in measuring the construct. In addition, since the values found by Hair et al. (2017) were appropriate for the acceptable range of $0.40 < x < 0.70$, no item removal was made.

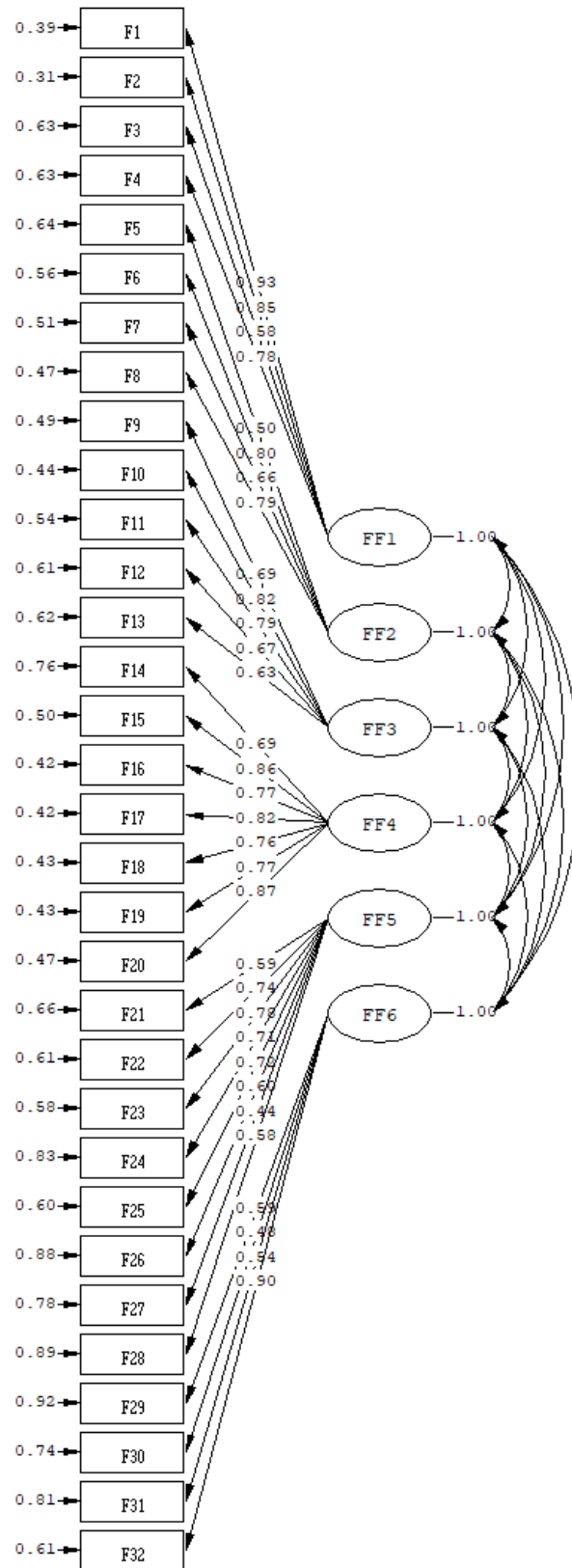


Figure 1. CFA Model

**** FF1: Discrimination, Exclusion and Toxicity, FF2: Information Hazards, FF3: Harms of Misinformation, FF4: Malicious Uses, FF5: Harms of Human-Computer Interaction, FF6: Automation, Access and Environmental Harms.**

Internal Consistency Analysis

The reliability of this scale was tested in terms of internal consistency with Cronbach α coefficient. The Cronbach α internal consistency coefficient of the 32 items in the scale was calculated as 0.932. For the factors, Cronbach's α internal consistency coefficient was 0.824, 0.764, 0.826, 0.898, 0.816, 0.700 respectively. These values are expected to be higher than 0.70 (Gefen, Straub, & Boudreau 2000; Hair, Anderson, Tatham, & Black, 1998). In this study, it can be said that the values provide sufficient evidence for the reliability of the scale.

Item-Total Scores Correlation

Item-total score correlations, which are used to express the relationship between the score for each item and the total score obtained from the scale, were calculated. It was determined that the item-total score correlation values ranged between 0.339 and 0.756 and were significant. A significant item-total score correlation indicates that the items have discrimination in terms of the measured feature (Büyüköztürk, 2004). In this case, it is seen that the total score correlations of the items are sufficient.

Item Discrimination Test

The total score obtained from the scale was ranked from high to low, and the scores of the group in the upper 27% group (N=80) and the group in the lower 27% group (N=80) were compared in terms of statistical differentiation. In the present study, a significant difference was observed between the upper and lower 27% groups according to the total test scores ($t=34.890$, $p=0.000$).

Structural Equation Model

Before testing the structural model, the validity and reliability of the measurement model was tested. Findings regarding the measurement model are presented in Supplementary Files 1. In this study, validity and reliability analyses of the measurement model including four constructs were conducted. First, the VIF values of the external model were examined and the fact that these values ranging between 1.102 and 2.908 were below 3 indicated that there was no multicollinearity problem.

The reliability of the constructs was assessed with Cronbach's alpha, combined reliability (CR), ρ_c and average variance explained (AVE) values. Reliability was above .70 for AI Literacy and Generative AI Risk Awareness constructs, while values close to .60 were observed for fixed and growth mindsets. Most of the AVE values approach the threshold value of .50 or are acceptable above .40 for theoretical reasons. Discriminant validity was tested with HTMT ratios and Fornell-Larcker criterion; HTMT values below .85 and AVE square root of each construct being higher than the correlations with other constructs revealed that the model provided sufficient discriminant validity. These findings indicate that the measurement model is satisfactory in terms of both convergent and discriminant validity.

After the necessary analyzes regarding the measurement model were made, the proposed structural model was tested with the Smart PLS 4.0 program. Structural model results are given in Figure 2.

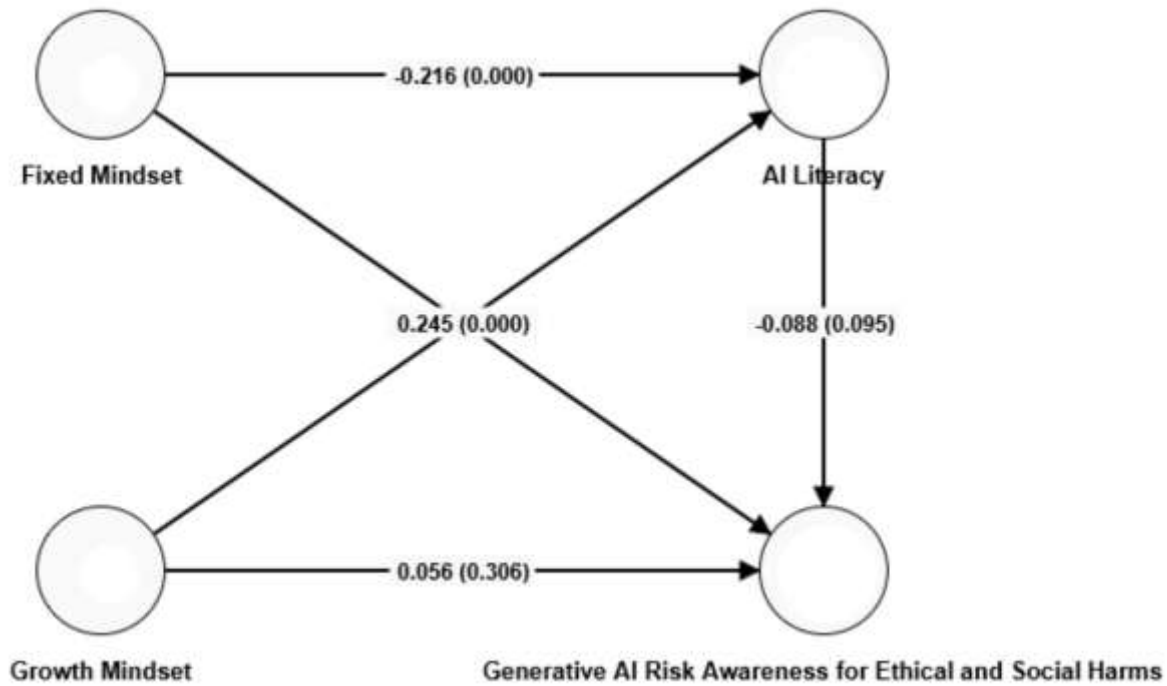


Figure 2. SEM Model

According to Figure 2 and Table 3, the hypothesis regarding the effect of GenAI literacy on GenAI risk Awareness for ethical and social harms was not supported ($\beta = -0.088$, $t = 1.670$, $p > .05$; H1 not supported). Fixed mindset's negative effect on GenAI literacy ($\beta = -0.216$, $t = 4.587$, $p < .01$; H2a accepted), and its positive effect on GenAI Risk Awareness for Ethical and Social Harms were confirmed ($\beta = 0.245$, $t = 4.939$, $p < .01$; H2b accepted). However, the effect of growth mindset on GenAI literacy ($\beta = -0.090$, $t = 1.676$, $p > .05$; H3a not supported) and GenAI Risk Awareness for Ethical and Social Harms was not statistically significant ($\beta = 0.056$, $t = 1.023$, $p > .05$; H3b not supported).

Table 3. Hypothesis Test Results

Hypothesis		Path	SD	T statistics	p
H1	GenAI Literacy -> GenAI Risk Awareness for Ethical and Social Harms	-0.088	0.053	1.670	0.095
H2a	Fixed Mindset -> GenAI Literacy	-0.216	0.047	4.587	0.000
H2b	Fixed Mindset -> GenAI Risk Awareness for Ethical and Social Harms	0.245	0.050	4.939	0.000
H3a	Growth Mindset -> GenAI Literacy	-0.090	0.054	1.676	0.094
H3b	Growth Mindset -> GenAI Risk Awareness for Ethical and Social Harms	0.056	0.055	1.023	0.306

Discussion

This research consists of two parts. The first one is scale development and adaptation and the other one is modeling. In this context, the GenAI ethics and social harm risk awareness scale was developed. Weidinger et al (2021) developed a 6-factor structure of GenAI ethics and social harm risk awareness with a comprehensive literature review. While the first three dimensions are related to ethics, the other 3 dimensions are related to social risks. The confirmatory factor analysis showed acceptable and perfect fit values for the fit indices and these values were reported to be significant. As a result of the appropriateness of the item loadings and expert opinion, no item deletion was made. In the 32-item scale, Cronbach α internal consistency coefficient was calculated as 0.932. For the factors, this coefficient was calculated as 0.824, 0.764, 0.826, 0.898, 0.816, 0.700 respectively. Since these values are higher than 0.70 (Gefen et al., 2000; Hair et al., 1998), it is concluded that there is sufficient evidence of reliability for the scale. Item-total score correlation values were measured for the discrimination level of the scale, and it was determined that these values ranged between 0.339 and 0.756 and were significant. In addition, it was seen that there was a significant difference between the 27% lower-upper groups according to the total test scores.

According to the results of confirmatory factor analysis for the GenAI literacy scale adapted into Turkish by integrating with the GenAI structure, it was seen that the fit indices were acceptable and had excellent fit. There are 4 dimensions and 10 items in the scale. The Cronbach α internal consistency coefficient for the 4 dimensions was calculated as 0.739. Cronbach α internal consistency coefficients for the sub-factors of the scale were 0.60, 0.682, 0.761, 0.60, respectively. It can be stated that these values of 0.60 and above provide evidence of reliability (Dijkstra & Henseler, 2015). It was determined that the item total score correlation values ranged between 0.300 and 0.656 and had a significant relationship. According to the test total scores, a significant difference was observed between the upper and lower 27% groups ($t=26.301$, $p=0.000$).

According to the results of the model analysis, GenAI literacy variable does not have a significant effect on GenAI risk awareness for ethical and social harms (H1, not accepted). The fact that the participants were not involved in a training program on GenAI ethical use and literacy may have led to this result. According to the results of an experimental study in which a holistic approach was used with AI problems with real life situation scenarios, it was stated that an AI literacy program improved AI ethic awareness (Kong et al., 2023). The inclusion of university students in activities for the use of this new technology in flexible learning activities as well as formal courses can differentiate this research result. GenAI literacy is beneficial for society at large as well as empowering individuals by taking advantage of ethical dilemmas and challenges. Public institutions and the private sector should be encouraged to take bottom-up (AI literacy, embedded ethical approaches) and top-down (regulations) measures. It is suggested that GenAI literacy should be conceptualized in the context of developer and user, and that the ethical principles underpinning AI-enabled mass customization and the universal challenges faced should be leveraged to propose ethically compatible mechanisms and policies (Hermann, 2022).

This article focuses on the impact of GenAI literacy on the level of awareness of developer-induced GenAI risk factors. At this point, educational institutions have important tasks in developing AI literacy and introducing

regulations. With the literacy courses to be included in the curricula related to GenAI, students can start learning with GenAI, otherwise they will only use GenAI to get information from GenAI (Chiu, 2023), and unethical use may increase, as in assignments given in the form of copying and pasting information from the internet. In order to adopt human-centered GenAI and be aware of the risks, there is a need to create the opportunity for more experience and literacy education about GenAI. Without the ability to understand how AI technology works, they will not be aware of the social and ethical risks, including being aware of its limitations and possibilities, and their ability to participate in a society in an equal and just way and to realize their own truths may be affected (Benton, 2023).

Fixed mindset has a significant effect on GenAI literacy (H2a, accepted). According to a study, it was concluded that individuals with a growth mindset react more positively to AI-supported smart technologies (Dang & Liu, 2022). Individuals with a fixed mindset give up easily in the face of obstacles, do not make much effort, and tend to avoid difficulties. Individuals with a growth mindset tend to embrace challenges, persevere in the face of obstacles, see effort as part of mastery, and learn from criticism (The Open University, 2023). GenAI has the potential to bring new pedagogical approaches to the agenda. Individuals with a fixed mindset may have developed literacy with the tool if they have gained a limited amount of experience of using GenAI in a way that supports its development, and if their experiences of use have been in the form of avoiding difficulties and obtaining information without effort. In other words, the way GenAI is used may be shaped by the lessons learned in this framework. If the ways of using GenAI that are supportive and positively affect (Gökçearslan et al., 2024) learning are encouraged, if policy makers clarify the rules on this issue and if these rules are implemented, individuals with a fixed mindset may not prefer the easy way, which is the first mindset.

Fixed mindset has a significant effect on GenAI risk awareness for ethical and social harms (H2b, accepted). It suggests that the person with a high fixed mindset would be less willing to make the decision to rely on the help of an AI to make a decision as there is a higher risk of performing poorly in risky situations (Salvesen and Møller, 2022). GenAI technology offers an environment where habits and patterns change. In a fixed mindset, change causes people to feel that their basic needs (social and physical connectedness) are threatened (Farrow, 2021). This may have led to the development of risk awareness. Individuals in a growth mindset are more adaptive, reaching higher levels and exceeding Maslow's (1943) basic needs. They have a higher level nature as if their basic needs are satisfied (Farrow, 2021). They seem to have overcome the security threat.

Growth mindset variable has no significant effect on GenAI literacy (H3a not accepted). Mindset plays an important role in the implementation of AI. However, a growth mindset is stated to be a necessary component to be successful in organizations and institutions driven by AI (Athota, 2021). Growth mindset does not have a significant effect on GenAI risk awareness for ethical and social harms (H3b, not accepted). According to Salvesen and Møller (2022), individuals with a growth mindset may be more adaptable to AI risks and less affected by them. It is noted that growth mindset is an important component of the future of literacy (Farrow, 2021). This research result may have been encountered due to insufficient exposure to risk factors and limited experience with GenAI, which may have resulted in individuals with a growth mindset not encountering an environment for learning and development.

Conclusion

The current study proposed a structural model showing the interrelations between AI literacy, GenAI awareness for ethical and social harm, and mindset. This model confirmed the negative impact of fixed mindset on AI literacy, while also finding its positive impact on ethical and social harms. As a result, conclusions were reached in terms of understanding the impact of individuals' thinking construct on Gen AI literacy and awareness of ethical issues. However, some limitations of the research should be considered. One of these is the limited understanding of the frequency and purpose and motivation of the participant group's Gen AI uses. Another limitation is the use of self-report surveys instead of using log data of students' actual usage.

Additionally, social norms and acceptances regarding the use of AI in the context of the research and the reflection of cultural characteristics on the findings should be taken into account. While the use of AI offers many potentials and facilitators for students, it also presents some uncertainties for students on issues such as assignment and project delivery. Students' understanding of ethical and unethical ways to use GenAI may influence their frequency of use and their views on AI. The relative newness of GenAI use may also influence student perceptions. In this context, the cross-sectional design of the current study can be seen as another limitation. Ensuring the continuity of studies on this subject through longitudinal studies can provide more insight into the robustness of research findings.

Implications, Recommendations and Future Research Directions

It is recommended to offer activities to improve the knowledge and practice of university students in the context of potential risks and ethical problems associated with GenAI technologies. It is recommended to offer activities to improve the knowledge and practice of university students in the context of potential risks and ethical problems associated with GenAI technologies. It is essential to organize workshops that incorporate practical activities, enabling students to use GenAI ethically. By designing these practice-based trainings to require students to make ethical decisions regarding GenAI usage, we can foster richer learning experiences. In addition, it is recommended that activities that encourage students in the social sciences to think about the social implications of AI should be included in the curriculum. In this context, students can be presented with cases to discuss GenAI's ethical and unethical uses. Individuals with different mindsets who develop behaviors in the face of challenges should be provided with customized learning opportunities in accordance with their mindsets in the first place. In the future, it is recommended that students should be directed to activities that support them to be in a growth mindset. Students can be supported to gain a more flexible approach with learning activities and workshop-based learning experiences that will make them aware of the role of fixed mindset tendencies in the approach to GenAI.

Faculty members, policy makers and students should collaborate together to continue their development on the potential harms and negative consequences of GenAI. In addition, it is recommended that decision-making authorities develop policies that clarify the limits, ethical rules, and sanctions for using AI in assignment submissions. Furthermore, it is imperative to create student guidelines that elucidate the ethical risks posed by generative AI and provide clear recommendations for its appropriate use (Esiyok et al., 2025). It is recommended

to include GenAI literacy in university curricula and encourage responsible GenAI use. Faculty members should be encouraged to integrate these technologies into their courses by improving their digital literacy levels in order to take advantage of AI opportunities in their course activities. The development of a generation of GenAI literate and ethically conscious individuals should be supported to face and overcome the complex problems arising from the rapid popularity of GenAI (Kamalov et al., 2023). In this context, it is recommended that higher education institutions provide faculty members with professional development program opportunities regarding AI integration in learning and teaching processes.

The developed and adapted measurement tools can be used to determine the GenAI literacy and risk levels of university students and it is useful to conduct research to improve these levels. In future research, the effect of mindset on GenAI literacy and GenAI ethical and social risk harm awareness can be revealed through experimental research. Future longitudinal studies are recommended to examine temporal changes in GenAI literacy and ethical and social risk harm. Additionally, studies analyzing student GenAI tool usage logs instead of self-reported measures could be considered a future research direction.

The model can be tested again with students with different levels of GenAI usage experience and different individual characteristics. The influence of various educational fields—such as social sciences, engineering, natural sciences, arts, and humanities—on the patterns and frequency of GenAI usage can be explored in future research. Additionally, cross-cultural studies are suggested to gain insights into the impact of cultural and social norms on this process. It is beneficial that the AI tools to be used are open source. We should be conscious when sharing our data with GenAI tools, pay attention to the confirmation of the information developed by the tools, update our awareness of risks, and ensure that the level of GenAI literacy is in line with the pace of development of this rapidly developing technology.

References

- Abuadas, M., & Albikawi, Z. (2025). Predicting nursing students' behavioral intentions to use AI: the interplay of ethical awareness, digital literacy, moral sensitivity, attitude, self-efficacy, anxiety, and social influence. *Journal of Human Behavior in the Social Environment*, 1-21. 10.1080/10911359.2025.2472852
- Akkaş, Ö. M., Tosun, C., & Gökçearslan, Ş., (2024). Artificial Intelligence (AI) and Cheating: The Concept of Generative Artificial Intelligence (GenAI). *Transforming Education With Generative AI* (pp.1-20), Pennsylvania: IGI Global.
- Athota, V. S. (2021). Why mindset matters in a digital age. *Mind Over Matter and Artificial Intelligence*, 1-15. https://doi.org/10.1007/978-981-16-0482-9_1
- Bahroun, Z. (2023). Transforming education: A comprehensive review of generative artificial intelligence in educational settings through bibliometric and content analysis. *Sustainability*, 15(17), 12983. <https://doi.org/10.3390/su151712983>
- Benton, P. (2023, November). AI Literacy: A Primary Good. In *Southern African Conference for Artificial Intelligence Research* (pp. 31-43). Cham: Springer Nature Switzerland.

- Chan, C. K. Y. (2023). A comprehensive AI policy education framework for university teaching and learning. *International Journal of Educational Technology in Higher education*, 20(1). 10.1186/s41239-023-00408-3
- Cheng, K. H., Liang, J. C., & Tsai, C. C. (2013). University students' online academic help seeking: The role of self-regulation and information commitments. *The Internet and Higher Education*, 16, 70-77. 10.1016/j.iheduc.2012.02.002
- Cheng, L., & Liu, X. (2023). From principles to practices: The intertextual interaction between AI ethical and legal discourses. *International Journal of legal Discourse*, 8(1), 31-52. 10.1515/ijld-2023-2001
- Chiu, T. K. (2023). The impact of Generative AI (GenAI) on practices, policies and research direction in education: A case of ChatGPT and Midjourney. *Interactive Learning Environments*, 1-17. 10.1080/10494820.2023.2253861
- Chiu, T. K. (2024). The impact of Generative AI (GenAI) on practices, policies and research direction in education: A case of ChatGPT and Midjourney. *Interactive Learning Environments*, 32(10), 6187-6203. <https://doi.org/10.1080/10494820.2023.2253861>
- Chugh, D., & Kern, M. C. (2016). Becoming as ethical as we think we are: The ethical learner at work. In *Organizational wrongdoing*. Cambridge University Press New York.
- Chui, M., Hazan, E., Roberts, R., Singla, A., & Smaje, K. (2023). The economic potential of generative AI. McKinsey Global Institute, June 2023
- Çelebi, C., Demir, U., & Karakuş, F. (2023). Examining studies on artificial intelligence literacy using the systematic review method. *Journal of Necmettin Erbakan University Ereğli Faculty of Education*, 5(2), 535-560.
- Çelik, H. E., & Yılmaz, V. (2013). LISREL 9.1 ile Yapısal Eşitlik Modellemesi: Temel Kavramlar-Uygulamalar-Programlama (2. Baskı). Ankara: Anı Yayıncılık.
- Çokluk, Ö., Şekercioğlu, G., & Büyüköztürk, Ş. (2012). *Sosyal bilimler İçin Çok Değişkenli İstatistik: SPSS ve LISREL Uygulamaları (2. Baskı)*. Ankara: Pegem Akademi.
- Dang, J., & Liu, L. (2022). A growth mindset about human minds promotes positive responses to intelligent technology. *Cognition*, 220, 104985. 10.1016/j.cognition.2021.104985
- Denny, P., Leinonen, J., Prather, J., Luxton-Reilly, A., Amarouche, T., Becker, B. A., & Reeves, B. N. (2023). Promptly: Using prompt problems to teach learners how to effectively utilize ai code generators. *arXiv preprint arXiv:2307.16364*.
- Dijkstra, T. K., & Henseler, J. (2015). Consistent partial least squares path modeling. *MIS Quarterly*, 39, 297-316. 10.25300/misq/2015/39.2.02
- Dweck, C. S. (2006). *Mindset: The new psychology of success*. Random house.
- Farrow, E. (2021). Mindset matters: How mindset affects the ability of staff to anticipate and adapt to Artificial Intelligence (AI) future scenarios in organisational settings. *AI & Society*, 36(3), 895-909. 10.1007/s00146-020-01101-z
- Esiyok, E., Gokcearslan, S., & Kucukergin, K. G. (2025). Acceptance of educational use of AI chatbots in the context of self-directed learning with technology and ICT self-efficacy of undergraduate students. *International Journal of Human-Computer Interaction*, 41(1), 641-650. <https://doi.org/10.1080/10447318.2024.2303557>

- Fui-Hoon Nah, F., Zheng, R., Cai, J., Siau, K., & Chen, L. (2023). Generative AI and ChatGPT: Applications, challenges, and AI-human collaboration. *Journal of Information Technology Case and Application Research*, 25(3), 277-304. 10.1080/15228053.2023.2233814
- Guo, K., Zhong, Y., Li, D., & Chu, S. K. W. (2023). Effects of chatbot-assisted in-class debates on students' argumentation skills and task motivation. *Computers & Education*, 203, 104862. 10.1016/j.compedu.2023.104862
- Gökcearslan, S., Tosun, C., & Erdemir, Z. G. (2024). Benefits, challenges, and methods of artificial intelligence (AI) chatbots in education: A systematic literature review. *International Journal of Technology in Education*, 7(1), 19-39. <https://doi.org/10.46328/ijte.600>
- Hair, J. F., Hult, G. T., Ringle, C. M., & Sarstedt, M. (2017). A primer on partial least squares structural equation modeling (PLS-SEM) (2nd ed.). Thousand Oaks, CA: Sage.
- Harrer, S. (2023). Attention is not all you need: The complicated case of ethically using large language models in healthcare and medicine. *EBioMedicine*, 90, 104512. <https://doi.org/10.1016/j.ebiom.2023.104512>
- Hermann, E. (2022). Artificial intelligence and mass personalization of communication content-An ethical and literacy perspective. *New Media & Society*, 24(5), 1258-1277. 10.1177/14614448211022702
- Hsieh, Y. H., & Tsai, C. C. (2014). Students' scientific epistemic beliefs, online evaluative standards, and online searching strategies for science information: The moderating role of cognitive load experience. *Journal of Science Education and Technology*, 23(3), 299-308. 10.1007/s10956-013-9464-6
- Kaltenegger, P. (2024). Fostering AI literacy: the mediating role of self-efficacy In artificial intelligence learning/submitted by Paul Kaltenegger, B. Sc. <https://epub.jku.at/obvulihs/content/titleinfo/11315302>
- Kamalov, F., Santandreu Calonge, D., & Gurrib, I. (2023). New era of artificial intelligence in education: Towards a sustainable multifaceted revolution. *Sustainability*, 15(16), 12451. 10.3390/su151612451
- Kane, K., Robinson-Combre, J., Zane, L., & Berge, Z. L. (2010). Tapping into social networking: Collaborating enhances both knowledge management and elearning. *Journal of Information and Knowledge Management Systems*, 40(1), 62-70. 10.1108/03055721011024928
- Kılıç, S. (2013). Örnekleme yöntemleri. *Journal of Mood Disorders*, 3(1), 44-46. <https://doi.org/10.5455/jmood.20130325011730>
- Kong, S. C., & Zhang, G. (2021). A conceptual framework for designing artificial intelligence literacy programmes for educated citizens. In *Conference proceedings (English paper) of the 25th Global Chinese Conference on Computers in Education (GCCCE 2021)* (pp. 11-15). Centre for Learning, Teaching and Technology, The Education University of Hong Kong.
- Kong, S. C., Cheung, W. M. Y., & Zhang, G. (2023). Evaluating an artificial intelligence literacy programme for developing university students' conceptual understanding, literacy, empowerment and ethical awareness. *Educational Technology & Society*, 26(1), 16-30. 10.1007/s10639-022-11408-7
- Lai, C. L. (2020). From organization to elaboration: Relationships between university students' online information searching experience and judgements. *Journal of Computers in Education*, 7(4), 463-485. 10.1007/s40692-020-00163-8
- Lee-Won, R. J., Joo, Y. K., Baek, Y. M., Hu, D., & Park, S. G. (2020). "Obsessed with retouching your selfies? check your mindset!": Female instagram users with a fixed mindset are at greater risk of disordered eating. *Personality and Individual Differences*, 167, 110223. 10.1016/j.paid.2020.110223

- Li, H., & Kim, S. (2024). Developing AI literacy in HRD: Competencies, approaches, and implications. *Human Resource Development International*, 27(3), 345-366. 10.1080/13678868.2024.2337962
- Lim, W. M., Gunasekara, A., Pallant, J. L., Pallant, J. I., & Pechenkina, E. (2023). Generative AI and the future of education: Ragnarök or reformation? A paradoxical perspective from management educators. *The International Journal of Management Education*, 21(2), 100790. 10.1016/j.ijme.2023.100790
- Lin, Z. (2023). Why and how to embrace AI such as ChatGPT in your academic life. *Royal Society Open Science*, 10(8). 10.1098/rsos.230658
- Májovský, M., Černý, M., Kasal, M., Komarc, M., & Netuka, D. (2023). Artificial intelligence can generate fraudulent but authentic-looking scientific medical articles: Pandora's box has been opened. *Journal of Medical Internet Research*, 25, e46924. 10.2196/46924
- Maslow A. H. (1943) A theory of human motivation. *Psychol Rev*, 50, 370-396
- Mason, R. O. (1986). Four ethical issues of the information age. *MIS Quarterly*, 10(1), 5-12.
- Mehan, V. ((2023). Exploring the future jobs, working experience, ethical issues and skills from artificial intelligence perspective. *International Journal of Innovative Science and Research Technology*, 8(9), 1144-1149. 10.5281/zenodo.8390544.
- Miao, F., & Holmes, W. (2023). Guidance for generative AI in education and research. <https://www.unesco.org/en/articles/guidance-generative-ai-education-and-research>.
- Mohamed, A. M. (2023). Exploring the potential of an AI-based Chatbot (ChatGPT) in enhancing English as a Foreign Language (EFL) teaching: Perceptions of EFL Faculty Members. *Education and Information Technologies*, 1-23. 10.1007/s10639-023-11917-z
- 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, 100041. 10.1016/j.caeai.2021.100041
- Ng, D. T., Luo, W., Chan, H. M., & Chu, S. K. (2022). Using digital story writing as a pedagogy to develop AI literacy among primary students. *Computers and Education: Artificial Intelligence*. <https://doi.org/10.1016/j.caeai.2022.100054>
- Niederman, F., & Baker, E. W. (2023). Ethics and AI Issues: Old Container with New Wine?. *Information Systems Frontiers*, 25(1), 9-28. 10.1007/s10796-022-10305-1
- Norman, M. K., Mayowski, C. A., & Fine, M. J. (2021). Authorship stories panel discussion: Fostering ethical authorship by cultivating a growth mindset. *Accountability in Research*, 28(2), 115-124. 10.1080/08989621.2020.1804374
- Nov, O., Singh, N., & Mann, D. (2023). Putting ChatGPT's medical advice to the (Turing) test: Survey study. *JMIR Medical Education*, 9, e46939. 10.2196/46939
- Panthier, C., & Gatinel, D. (2023). Success of ChatGPT, an AI language model, in taking the French language version of the European Board of Ophthalmology examination: A novel approach to medical knowledge assessment. *Journal Français d'Ophtalmologie*, 46(7), 706-711. 10.1016/j.jfo.2023.05.006
- Pinski, M., & Benlian, A. (2023). Ai literacy-towards measuring human competency in artificial intelligence. 56th Hawaii International Conference on System Sciences. Hawaii. <https://hdl.handle.net/10125/102649>
- Rozak, A., & Karman, K. (2025). Digital literacy: Understanding norms and ethics in the online realm. *International Journal Of Humanities, Social Sciences And Business (INJOSS)*, 4(2), 36-42.
- Salvesen, H. B., & Møller, H. T. (2022). Trusting Artificial Intelligence: How Perceived Situational Risik and

- Digital Mindset Can Influence Individuals Decision to Rely (Master's thesis, Handelshøyskolen BI). <https://hdl.handle.net/11250/3034746>
- Schumacker, R. E., & Lomax, R. G. (2004). *Beginner's Guide to Structural Equation Modeling* (2nd Ed.). New Jersey: Lawrence Erlbaum Associates
- Sison, A. J. G., Daza, M. T., Gozalo-Brizuela, R., & Garrido-Merchán, E. C. (2023). ChatGPT: More than a weapon of mass deception, ethical challenges and responses from the human-Centered artificial intelligence (HCAI) perspective. *arXiv preprint arXiv:2304.11215*.
- Strauß, S. (2021). "Don't let me be misunderstood": Critical AI literacy for the constructive use of AI technology. *TATuP-Zeitschrift für Technikfolgenabschätzung in Theorie und Praxis/Journal for Technology Assessment in Theory and Practice*, 30(3), 44-49. 10.14512/tatup.30.3.44
- Strauß, S. (2021). "Don't let me be misunderstood": Critical AI literacy for the constructive use of AI technology. *TATuP-Zeitschrift für Technikfolgenabschätzung in Theorie und Praxis/Journal for Technology Assessment in Theory and Practice*, 30(3), 44-49. <https://doi.org/10.14512/tatup.30.3.44>
- Su, Y., Lin, Y., & Lai, C. (2023). Collaborating with ChatGPT in argumentative writing classrooms. *Assessing Writing*, 57, 100752. 10.1016/j.asw.2023.100752
- The Open University (2023). Summary of Dweck's Theory. 01.12.2023 tarihinde web üzerinde: <https://www.open.edu/openlearn/mod/oucontent/view.php?id=65520§ion=2>
- Tirado-Olivares, S., Navío-Inglés, M., O'Connor-Jiménez, P., & Cózar-Gutiérrez, R. (2023). From human to machine: Investigating the effectiveness of the conversational AI ChatGPT in historical thinking. *Education Sciences*, 13(8), 803. 10.3390/educsci13080803
- Tlili, A., Shehata, B., Adarkwah, M. A., Bozkurt, A., Hickey, D. T., Huang, R., & Agyemang, B. (2023). What if the devil is my guardian angel: ChatGPT as a case study of using chatbots in education. *Smart Learning Environments*, 10(1), 15. 10.1186/s40561-023-00237-x
- Tsai, C. C. (2009). Conceptions of learning versus conceptions of web-based learning: The differences revealed by college students. *Computers & Education*, 53(4), 1092-1103. 10.1016/j.compedu.2009.05.019
- Tseng, S. C., Liang, J. C., & Tsai, C. C. (2014). Students' self-regulated learning, online information evaluative standards and online academic searching strategies. *Australasian Journal of Educational Technology*, 30(1). 10.14742/ajet.242
- Wach, K., Duong, C. D., Ejdys, J., Kazlauskaitė, R., Korzynski, P., Mazurek, G., Paliszkievicz, J. & Ziemba, E. (2023). The dark side of generative artificial intelligence: A critical analysis of controversies and risks of ChatGPT. *Entrepreneurial Business and Economics Review*, 11(2), 7-30. 10.15678/eber.2023.110201
- Wang, B., Rau, P. L. P., & Yuan, T. (2023). Measuring user competence in using artificial intelligence: validity and reliability of artificial intelligence literacy scale. *Behaviour & information technology*, 42(9), 1324-1337. 10.1080/0144929X.2022.2072768
- Wang, B., Patrick Rau, P., & Yuan, T. (2023). Measuring user competence in using artificial intelligence: Validity and reliability of artificial intelligence literacy scale. *Behaviour & Information Technology*, 42(9), 1324-1337. 10.1080/0144929X.2022.2072768
- Wang, X., Liu, Q., Pang, H., Tan, S. C., Lei, J., Wallace, M. P., & Li, L. (2022). What matters in AI-supported learning: A study of human interactions using cluster analysis and epistemic network analysis. *Computers & Education*, 194, 104703. 10.1016/j.compedu.2022.104703

- Weidinger, L., Mellor, J., Rauh, M., Griffin, C., Uesato, J., Huang, P. S., ... & Gabriel, I. (2021). Ethical and social risks of harm from language models. *arXiv preprint arXiv:2112.04359*.
- Weinberg, L. (2022). Rethinking fairness: An interdisciplinary survey of critiques of hegemonic ML fairness approaches. *Journal of Artificial Intelligence Research*, 74, 75-109. 10.1613/jair.1.13196
- Wu, Y. T., & Tsai, C. C. (2005). Information commitments: Evaluative standards and information searching strategies in web-based learning environments. *Journal of Computer Assisted Learning*, 21(5), 374-385. 10.1111/j.1365-2729.2005.00144.x
- Yilmaz, E. (2022). Development of mindset theory scale (growth and fixed mindset): a validity and reliability study (Turkish version). *Research on Education and Psychology*, 6(Special Issue), 1-26. 10.54535/rep.10542

Author Information

Sahin Gokcearslan

 <https://orcid.org/0000-0003-3532-4251>.

Hacettepe University

Institute of Informatics

Beytepe, Ankara

Türkiye

Contact e-mail: sahingokcearslan@hacettepe.edu.tr

Hatice Yildiz Durak

 <https://orcid.org/0000-0002-5689-1805>

Necmettin Erbakan University

Department of Educational Science, Instructional Technology

Türkiye

Mustafa Serkan Gunbatar

 <https://orcid.org/0000-0003-0485-3038>

Van Yuzuncuyil University

Faculty of Education

Department of Computer Education and Instructional Technology

Türkiye

Nilufer Atman Uslu

 <https://orcid.org/0000-0003-2322-4210>

Dokuz Eylul University

Buca Faculty of Education

Department of Elementary Education

Türkiye

Aysun Nuket Elci

 <https://orcid.org/0000-0002-0200-668X>

Dokuz Eylul University

Buca Faculty of Education

Department of Elementary Education

Türkiye

Appendix A. GenAI Ethics and Social Risk Awareness Scale for University Students [Turkish Version]

Üniversite Öğrencileri için GenAI Etiği ve Sosyal Risk Farkındalığı Ölçeği

Ayrımcılık, Dışlama ve Toksisite: F1

1. Üretken yapay zeka belirli gruplara ayrımcılık yapar.
2. Üretken yapay zeka belirli grupları dışlar.
3. Üretken yapay zeka toksik bir dil (zarar verici, çatışmacı vb.) sergiler.
4. Üretken yapay zeka bazı sosyal gruplara düşük performans sergiler.

Bilgi Tehlikeleri: F2

5. Üretken yapay zeka yanlış veya yanıltıcı bilgi üretir.
6. Üretken yapay zeka kişisel bilgilerin gizliliğini tehlikeye atar.
7. Üretken yapay zeka siber saldırılara karşı savunmasız olabileceği için yanlış bilgi üretir.
8. Üretken yapay zeka hassas bilgileri sızdırır.

Yanlış Bilginin Zararları: F3

9. Üretken yapay zeka yanlış veya yanıltıcı bilgi üreterek zarara neden olur.
10. Üretken yapay zeka sağlıkla ilgili yanlış veya yanıltıcı bilgileri yayarak maddi zarara neden olur.
11. Üretken yapay zeka sağlıkla ilgili yanlış veya yanıltıcı bilgileri yayarak sağlık problemlerine neden olur.
12. Üretken yapay zeka kullanıcıları etik olmayan eylemlerde bulunmaya teşvik eder.
13. Üretken yapay zeka kullanıcıları yasa dışı eylemlerde bulunmaya teşvik eder.

Kötü Niyetli Kullanımlar: F4

14. Üretken yapay zeka kötü niyetli insanlar tarafından manipülasyon amacıyla kullanılır.
15. Üretken yapay zeka siber saldırı veya veri hırsızlığı gibi kötü niyetli kullanıma hizmet eder.
16. Üretken yapay zeka ile kişisel bilgiler kötü niyetli kullanılır.
17. Üretken yapay zeka yasa dışı faaliyetler için kullanılır.
18. Üretken yapay zeka insan hayatını tehlikeye atabilecek biçimde kullanılır.
19. Üretken yapay zeka sağlığı tehlikeye atacak şekilde insanları manipüle eder.
20. Üretken yapay zeka insan hak ve özgürlüklerini ihlal edebilecek biçimde kullanılır.

İnsan-Bilgisayar Etkileşiminin Zararları: F5

21. Üretken yapay zeka kullanıcılarda hüsrana yol açar .
22. Üretken yapay zeka kullanıcılarda kafa karışıklığına neden olur.
23. Üretken yapay zeka kullanıcılarda endişe veya strese neden olur.
24. Üretken yapay zeka kullanıcılarda sosyal izolasyona (etkileşimin ve kişisel ilişkilerin azalması) neden olur.
25. Üretken yapay zeka kullanıcılarda etik ikilemlere neden olur.
26. Üretken yapay zeka kullanıcıların kişisel bilgilerine erişir.
27. Üretken yapay zekanın oluşturduğu bilgilere fazla güvenmek kullanıcıların karar verme süreçlerini etkiler.
28. Üretken yapay zeka cinsiyet ve etnik kimlik gibi özellikleri ima ederek kullanıcılarda önyargı oluşturur.

Otomasyon, Erişim ve Çevresel Zararlar: F6

29. Üretken yapay zekanın insan emeğine dayalı iş gücünün yerini alır.

30. Üretken yapay zeka teknolojiye erişim imkanı farklı olan insanlar arasındaki uçurumu artırır.
 31. Üretken yapay zeka çevreye zarar verir.
 32. Üretken yapay zeka insan emeğine dayalı bazı ekonomilere zarar verir.
-

Appendix B. GenAI Ethics and Social Risk Awareness Scale for University Students

[English Version]

GenAI Ethics and Social Risk Awareness Scale for University Students

Discrimination, Exclusion and Toxicity: F1

1. GenAI discriminates against certain groups.
 2. GenAI excludes certain groups.
 3. GenAI exhibits toxic language (damaging, confrontational, etc.).
 4. GenAI exhibits low performance to some social groups.
-

Information Hazards: F2

5. GenAI generates false or misleading information.
 6. GenAI compromises the privacy of personal information.
 7. GenAI generates false information because it could be vulnerable to cyber-attacks.
 8. GenAI leaks sensitive information.
-

The Harm of Misinformation: F3

9. GenAI causes harm by generating false or misleading information.
 10. GenAI causes financial harm by spreading false or misleading health-related information.
 11. GenAI causes health problems by spreading false or misleading health-related information.
 12. GenAI encourages users to engage in unethical actions.
 13. GenAI encourages users to commit illegal acts.
-

Malicious Uses: F4

14. GenAI is used for manipulation by people with malicious intent.
 15. GenAI serves malicious uses, such as cyber-attacks or data theft.
 16. Personal information is used maliciously with GenAI.
 17. GenAI is used for illegal activities.
 18. GenAI is used in a way that could endanger human life.
 19. GenAI manipulates people in a way that endangers health.
 20. GenAI is used in a way that may violate human rights and freedoms.
-

Harms of Human-Computer Interaction: F5

21. GenAI leads to user frustration.
 22. GenAI causes confusion in users.
 23. GenAI causes anxiety or stress in users.
 24. GenAI causes social isolation (reduced interaction and personal relationships) in users.
 25. GenAI causes ethical dilemmas in users.
 26. GenAI accesses users' personal information.
 27. Relying too much on the information generated by GenAI affects users' decision-making processes.
-

28. GenAI creates bias in users by implying characteristics such as gender and ethnic identity.
-

Automation, Access and Environmental Damages: F6

-
29. GenAI replaces the reliance on human labor.
 30. GenAI increases the gap between people with different access to technology.
 31. GenAI harms the environment.
 32. GenAI harms some economies based on human labor.
-