



Research Article

Revisiting the interplay of risks, expected benefits, and ethical implications of AI across varying automation and criticality levels in Arab and UK contexts

Mohammad Mominur Rahman^{a,*}, Areej Babiker^a, Ala Yankouskaya^b, Sameha AlShakhsi^a, Raian Ali^{a,*}

^a College of Science and Engineering, Hamad Bin Khalifa University, Doha, Qatar

^b School of Psychology, Bournemouth University, Bournemouth, UK

ARTICLE INFO

Keywords:

Artificial Intelligence
Automation
Criticality
Risks
Positive change
Ethical implications

ABSTRACT

Understanding AI's ethical implications is vital for its responsible deployment and societal acceptance. This study revisits the relationship between AI's potential risks, the positive changes it brings, and its ethical evaluations across cultural contexts, focusing on different levels of automation and criticality. An online survey involving 323 participants from the Arab GCC and 316 from the UK assessed responses to four scenarios, varying by low/high criticality and low/high automation. Hierarchical linear regression examined how perceptions of risk and positive change influence ethical evaluations. Findings show that perceived risks consistently raise ethical concerns, while positive changes mitigate these concerns, especially in high-risk scenarios. Although effects vary between UK and Arab participants, positive changes moderated the relationship between risks and ethical implications in both groups under high criticality. The study provides empirical evidence that perceived risks and benefits jointly influence ethical evaluations across different AI modalities, while cross-cultural variation remains modest.

1. Introduction

Artificial Intelligence (AI) technologies have become ubiquitous, profoundly impacting various aspects of our society and daily lives. From predictive algorithms aiding healthcare diagnostics to the development of autonomous vehicles and personalized learning systems, AI's reach is extensive and continuously evolving (Kaswan et al., 2024; Ma et al., 2020). AI facilitates automation, enhances decision-making, and drives innovation, making it a key enabler in the knowledge-based economy. Its potential to create new business models for small and medium sized enterprises (SMEs) (von Garrel & Jahn, 2023) and improve supply chain resilience in Industry 5.0 highlights its transformative impact across sectors (Wu et al., 2024). While offering immense potential for innovation and efficiency, the integration of AI into societal norms and ethical frameworks raises complex questions. On the one hand, AI promises to enhance human life by facilitating intercultural collaboration, guiding online consumer decisions, and streamlining tedious tasks (Davenport & Ronanki, 2018; Kamilaris & Prenafeta-Boldú, 2018). Conversely, concerns about AI's potential to

displace human jobs, facilitate the development of intelligent weaponry, and erode control over emerging technologies abound. The emergence of AI-driven text and image-generation tools has sparked debates surrounding artistic integrity, academic integrity, and the intrinsic value of human creativity (T. Brown et al., 2020; Elgammal et al., 2017). This dichotomy underscores the multifaceted nature of the AI discourse, necessitating a comprehensive examination of its ethical, legal, and social implications.

Ethical considerations in AI encompass a wide array of concerns, including privacy, accountability, fairness, transparency, and the potential for societal disruption. Research indicates that factors such as cultural background, personal values, and prior experiences shape individuals' perceptions of AI ethics. Responsible technology emphasizes the design, deployment, and use of technological systems in ways that promote human flourishing and societal well-being. It seeks to address ethical challenges such as privacy, fairness, transparency, and accountability, while fostering sustainable innovation that aligns with societal values (Stilgoe et al., 2013). These considerations underscore the importance of addressing significant ethical concerns, such as

* Corresponding authors.

E-mail addresses: mora28982@hbku.edu.qa (M.M. Rahman), arbabiker@hbku.edu.qa (A. Babiker), ayankouskaya@bournemouth.ac.uk (A. Yankouskaya), salshakhsi@hbku.edu.qa (S. AlShakhsi), raali2@hbku.edu.qa (R. Ali).

<https://doi.org/10.1016/j.jrt.2026.100163>

algorithmic biases and privacy breaches, through mitigation strategies that ensure AI technologies promote human flourishing. This involves embedding societal values into AI systems to safeguard public interests and uphold fundamental rights (Stahl, 2021). For instance, studies have shown that cultural norms and societal values significantly influence people's attitudes toward privacy and data protection in AI systems (Huriye, 2023). Additionally, individuals with a stronger commitment to ethical principles may be more likely to scrutinize AI systems for potential biases or unfairness (Kaswan et al., 2024). Moreover, psychological factors such as trust in AI systems and perceived control over their use can also impact ethical evaluations. Research suggests that individuals are more likely to accept AI technologies when they perceive them as trustworthy and transparent (Wanner et al., 2022). Similarly, feelings of control over AI systems may lead to more favorable ethical assessments, as people may feel empowered to mitigate risks (Figueras et al., 2021). In healthcare, AI ethics emphasizes transparency, data privacy, and the protection of patient autonomy, with core challenges including the need for explainable AI (XAI) methods and the alignment of system design with user values to build trust while preventing manipulation and overreliance (Calvaresi et al., 2025). Recently, authors argue that inconsistent ethical regulation of research with emerging digital technologies, including AI, may undermine protections for individuals and society, stressing the need for harmonized standards (Mollen, 2024). These considerations emphasize that embedding societal values into AI systems is essential for addressing ethical concerns, mitigating risks, and ensuring responsible deployment across diverse cultural contexts.

Besides personality factors, studies have shown that individuals' perceptions of the risks associated with AI technologies are influenced by factors such as media coverage, personal experiences, and trust in institutions. Media coverage highlighting ethical issues of AI technologies and including accessibility of correct information to the public can lead to greater scrutiny of their ethical implications (Ouchchy et al., 2020). Negative media portrayals of AI, highlighting potential job displacement or privacy breaches, can lead to heightened risk perceptions and ethical concerns (Brossard & Scheufele, 2013). Additionally, individuals with past negative experiences or lack of trust in AI developers may be more inclined to perceive AI technologies as risky and ethically questionable (Glikson & Woolley, 2020).

Individuals' perceptions of the positive changes enabled by AI technologies can also shape their ethical evaluations. For instance, AI's potential to improve healthcare outcomes, enhance productivity, and drive innovation may lead to more favorable ethical assessments (Mulgan, 2016). Positive personal experiences with AI applications, such as virtual assistants or recommendation systems, can foster trust and acceptance, influencing ethical evaluations (Xiong et al., 2023). Importantly, individuals engage in a cognitive weighing process when evaluating the ethics of AI technologies, considering both their perceived risks and benefits (Kamila & Jasrotia, 2023). This balancing act between potential harms and advantages plays a crucial role in shaping ethical assessments and guiding decision-making regarding AI adoption and regulation.

Behavioral economics offers valuable insights into understanding public responses to AI. Camerer's study explores how AI and behavioral economics intersect in three key ways: using machine learning (ML) to predict choices, modeling human judgment as implicit ML prone to overfitting, and how AI can assist preference formation while also enhancing price discrimination (Camerer, 2019). Researcher further highlights how integrating ML with behavioral economics can enhance decision-making through personalized interventions, stronger predictive power, and more effective policy design (Hrnjic & Tomczak, 2019). By applying these behavioral economics principles, policymakers and AI developers can better address public concerns, such as risk aversion and trust issues, particularly in high-risk AI applications such as autonomous vehicles, where complacency can skew risk perception (Azuma et al., 2023), and healthcare systems, where trust in AI must be balanced with cautious decision-making by clinicians (Choudhury, 2022).

The mode of AI operation also matters. The degree of automation and criticality of AI technologies plays a significant role in shaping individuals' perceptions of their ethical implications. Automation refers to the extent to which a system operates without human intervention, a concept closely linked to perceptions of user control, accountability, and trust (Parasuraman et al., 2000). Research in moral psychology shows that people's aversion to automated moral decision-making stems in part from a perceived loss of human oversight, particularly in high-stakes contexts, where machines are seen as lacking moral agency (Bigman & Gray, 2018). Criticality, defined as the severity of consequences resulting from system failures, is associated with heightened ethical scrutiny, as moral judgments are strongly influenced by emotional engagement with perceived harm (Greene et al., 2001). Empirical studies further demonstrate that individuals respond more strongly, both emotionally and ethically, to AI systems deployed in high-criticality domains such as healthcare or autonomous driving, compared to lower-stakes settings like cleaning or entertainment (Awad et al., 2018a). Research suggests that the level of automation, or the extent to which AI systems operate autonomously, and the criticality of tasks they perform influence how people evaluate the risks and benefits of AI technologies (Ajenaghughure et al., 2020; Ward et al., 2017). For instance, in scenarios where AI applications exhibit low levels of automation and are used in low-criticality tasks, individuals may perceive fewer risks and more positive changes associated with these technologies (Awad et al., 2018b). Conversely, in situations involving high levels of automation and high-criticality tasks, such as in autonomous vehicles or medical diagnosis systems (Morley et al., 2020), individuals may have heightened concerns about risks and ethical considerations (Banks et al., 2019). Studies have shown that perceptions of automation and criticality interact with other factors, such as trust in AI systems and personal experiences, to shape individuals' ethical evaluations. For example, individuals may perceive greater risk associated with an AI system if they have had negative experiences using it, which will impact their trust in the system regardless of the system type or task (Morley et al., 2020). Conversely, positive experiences or a sense of trust in AI systems may lead individuals to perceive more positive changes and benefits, particularly in low-criticality contexts with lower levels of automation (Park & Jung, 2021). Moreover, automation also has the potential to streamline administrative processes, reduce manual errors, and free up valuable resources for more critical tasks (Parycek et al., 2024). These dynamics are consistent with the IMPACT model (Montag et al., 2024), which emphasizes the interplay between individual characteristics, modality, context, cultural/country-level influences, and transparency as critical components in shaping public attitudes toward AI.

This study examines how AI's degree of autonomy and task criticality impacts ethical evaluations. Several studies look at these factors separately. Klein et al. (2024) tested autonomy levels at different impact levels and found that both influenced trust and risk perception. Tolmeijer et al. (2022) studied autonomy in ethical decision-making and discovered it affected perceived responsibility. Kieslich, et al. (2022) surveyed ethical principles related to machine autonomy and identified different preference groups. However, Chanseau et al. (2017) is the only source that directly explores the link between task criticality and autonomy preferences, although it focused on domestic robots instead of ethical evaluations in a wider sense. Thus, while these factors are conceptually important in ethical debates, prior research has offered limited systematic empirical comparison of how they jointly shape public perceptions across different operational contexts. By incorporating these factors into the analysis, the research provides a more granular view of how AI's operational context affects public perception and acceptance. This study builds on prior work by systematically examining whether the perceived benefits of AI ("positive changes"), moderate ethical concerns in autonomy and criticality modalities across two cultural contexts. This aspect is often underexplored in ethical evaluations of AI technologies. By considering how positive outcomes, can offset ethical apprehensions, the research sheds light on the

balancing act individuals engages in when evaluating AI. This approach contributes to a more comprehensive understanding of how to navigate AI's ethical challenges, particularly in high-autonomy, high-criticality scenarios. This study extends trust and acceptance models by focusing on ethical implications as a distinct evaluative outcome and examining how perceived risks and benefits jointly shape these judgments under systematically varied AI conditions. Extending prior research on cultural variation in AI perceptions, where application type is often specified only at a general level, we manipulate task criticality and autonomy in a 2×2 factorial design. This approach enables within-mode and between-mode comparisons, offering a more precise understanding of how operational context interacts with cultural factors to shape ethical evaluations of AI. Furthermore, our study explores the role of cultural nuances in shaping individual propensities for concerns and intentions related to AI use. Recent studies have emphasized the overrepresentation of WEIRD (Western, Educated, Industrialized, Rich, and Democratic) contexts in AI ethics, HCI, and learning analytics research; e.g., 84 % of ACM FAccT studies involved Western participants (Septiandri et al., 2023), 73 % in CHI (Linxen et al., 2021), and 58 % in learning analytics (Baek & Doleck, 2024). Using data from the UK and an Arab sample, this study contributes cross-cultural evidence on AI perceptions and suggests that many patterns hold across contexts. In AI-specific contexts, recent findings from Rahman et al. (Rahman et al., 2024) show that demographic factors influence AI acceptance differently in the UK versus the Arab, underscoring the value of continued cross-cultural investigation.

The coming wave of AI technology, as described by Mustafa Suleyman in his book (Suleyman, 2023), is unstoppable and demands concerted efforts to harness its benefits while minimizing and mitigating the ethical implications influenced by its potential risks. Understanding how these factors intersect and influence individuals' ethical evaluations of AI technologies is essential for policymakers, developers, and stakeholders. By identifying the underlying drivers of ethical considerations, policymakers can design regulations and governance models that address public concerns and promote responsible AI development (Huriye, 2023). Similarly, AI developers might be mandated to enhance the explainability of risks and benefits in a way that aids a calibrated ethical judgment. This entails changes in the design and deployment of AI systems, ensuring that they align with societal values and expectations (Dignum, 2017).

The present study addresses the following research questions across UK and Arab samples:

1. For AI operating under different autonomy and criticality conditions, how do perceptions of potential risks influence perceptions of ethical concerns?

2. For AI operating under different autonomy and criticality conditions, what role do perceived positive changes brought by the technology play in moderating the relationship between perceptions of its potential risks and ethical concerns?

Based on these questions, the study proposes the following hypotheses across the different configurations of AI considering autonomy and criticality levels, as demonstrated in Fig. 1. These hypotheses are informed by two key theories in moral psychology: the dual-attitude model (Wilson et al., 2000) and Cushman's dual-system framework (Cushman, 2013). From the dual-attitude model, we speculate that implicit evaluations may influence risk sensitivity, whereas explicit evaluations may more strongly reflect benefit assessments, such as AI. This helps explain why individuals may perceive AI as risky while also recognizing its potential benefits. Cushman's framework adds that moral judgments often integrate an action-based system (often linked to affective harm aversion) and an outcome-based system (grounded in deliberative evaluation of consequences). Together, these perspectives support our expectation that perceived risks heighten ethical concerns, while perceived benefits may mitigate or moderate those concerns.

H1: Higher potential risks of AI technologies are associated with greater ethical concerns, while the positive changes they bring tend to reduce these concerns.

H2: Positive changes brought by AI technologies moderate and lessen the negative influence that potential risks have on ethical concerns.

2. Method

2.1. Participant recruitment and selection

The participant pool for this study comprised 639 adults, comprising 323 participants from the Arab GCC and 316 from the UK, according to their nationality, representing a diverse cultural spectrum. The required sample size was estimated using G*Power (Bruin, 2011), applying the F-test family for Linear Multiple Regression: Fixed Model, R^2 increase from zero, with 13 predictors. Based on a small-to-medium effect size ($f^2 = 0.08$), a significance level of $\alpha = 0.05$, and a desired power of 0.80, the analysis indicated a minimum total sample of $N = 235$. For group comparisons, an additional power analysis was conducted using the *t*-test family (two-tailed test for independent means), with an allocation ratio of 1.022 (Arab sample to UK sample), $\alpha = 0.05$, power = 0.90, and a minimum detectable effect size of Cohen's $d = 0.30$. This yielded a required sample of 232 for the Arab group and 238 for the UK group, both thresholds were exceeded in the final dataset.

While our UK sample did include responses from older individuals,

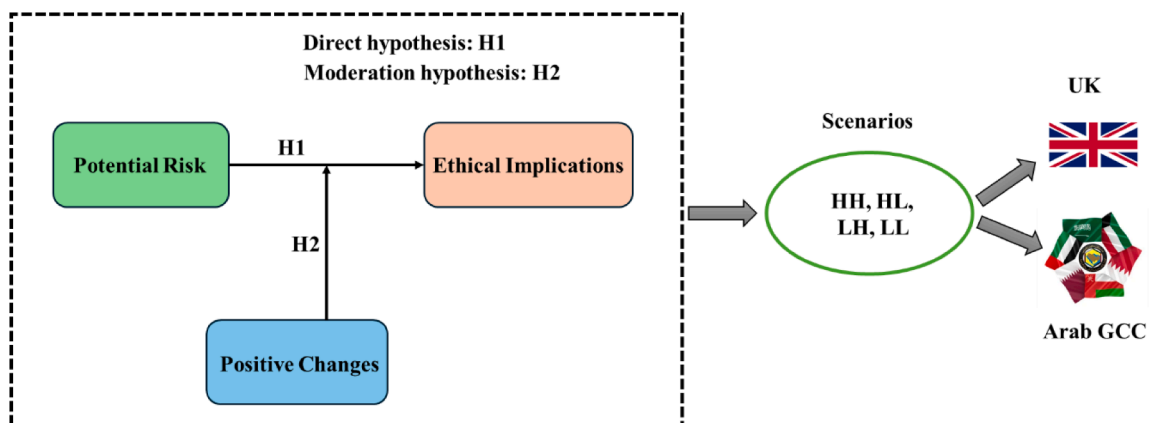


Fig. 1. Conceptual model for hierarchical linear regression analysis. This model will be applied to different combinations of autonomy (low, high) and criticality (low, high) across both UK and Arab samples. L: Low, H: High.

we opted to restrict the sample to ages between 18 and 60 years to address the challenge of recruiting Arabs above this age. Participants were required to complete a preliminary screening questionnaire to qualify for participation in the main survey. To be eligible, participants needed to be 18 years or older, native to and residing in either the UK or the Arab GCC, possess some familiarity with AI, and self-identify culturally as either British or Arab. All participants were given informed consent forms and were allowed to withdraw from the survey at any point. To maintain data quality, the survey included attention checks. Participants who failed these checks or completed the survey in less than half of the median time were excluded from the analysis.

Compensation was provided to participants upon completion of the survey. Ethical clearance was secured from the Institutional Review Board (IRB) of the lead author's institution, and participants gave informed consent, with the option to withdraw from the survey at any stage. The survey was administered online via the SurveyMonkey platform (surveymonkey.com) (Waclawski, 2012), spanning from late October 2023 to mid-December 2023. TGM Research, a market research firm with an expansive online panel presence across 130 countries, facilitated recruitment (Greg Laski, n.d.). Participation was incentivized through TGM Research, adhering to the ethical standards of voluntary participation and confidentiality. This study, part of a larger project, will concentrate on the sections that directly pertain to the participants' perceptions of risks, positive changes, and ethical concerns of four AI scenarios representing different modes of operation. The dataset utilized in this study is available at the Open Science Framework (OSF) link provided in the Data Availability Statements section.

2.2. Measures

Consistent with experimental survey practices, this study utilized vignette experiments to examine public attitudes toward AI. Vignette experiments involve presenting respondents with brief, scenario-based descriptions that prompt evaluative judgments. This approach captures more detailed and context-sensitive responses compared to abstract questioning (Alexander & Becker, 1978). It also enhances the reliability and validity of responses by enabling controlled variation of scenario elements through methods like confounded factorial designs, between-subjects manipulations, and the use of anchoring vignettes (Steiner et al., 2016). Four vignettes were created and face-validated, varying in terms of AI's automation and criticality levels. Each vignette varied in two theoretically important dimensions, AI autonomy and task criticality, enabling a factorial structure that supports the analysis of main and interaction effects. Participants were randomly assigned to evaluate two of the four vignettes, ensuring a balanced and interpretable design while minimizing cognitive fatigue. This assignment strategy aligns with recommendations for subset distribution in factorial vignette designs, allowing for meaningful comparisons while reducing respondent burden (Atzmüller & Steiner, 2010). Together, this approach allowed us to systematically explore how scenario characteristics influence ethical evaluations across cultural contexts. Participants were introduced to the concept of automation and criticality levels through the following explanations:

"The following section of the survey will introduce AI technology characteristics followed by scenarios of using AI in everyday situations. We will primarily utilize two AI-powered systems: a cleaning robot and a smart car.

Automation Level: Describes AI's autonomy in decision-making and action-taking. High automation indicates AI's self-sufficiency, while low automation requires more human guidance.

Criticality Level: Measures AI's impact on humans. High criticality implies serious consequences from AI errors, while low criticality suggests manageable effects".

following statements to measure potential risks, positive changes, and ethical implications, respectively:

- 1) "I am concerned about potential risks associated with this AI technology." The degree of risk was measured on a 6-point Likert scale ranging from "1 = Strongly disagree" to "6 = Strongly agree".
- 2) "I believe that this AI application can bring about positive changes, both personally and for society." This was also quantified using a 6-point Likert scale from "1 = Strongly disagree" to "6 = Strongly agree".
- 3) "To what extent are you concerned with the ethical implications of this AI technology?" Respondents could indicate their level of concern on a more nuanced 11-point scale, with endpoints labeled as 0 for "I am not at all concerned" and 10 for "I am completely concerned."

These measures were newly constructed for this study to align with the specific context of AI automation and criticality levels in a vignette-based design, as existing validated scales were not sufficiently tailored to our factorial design. The use of single-item measures was chosen to reduce participant fatigue, given that participants evaluated multiple scenarios, and is supported by evidence that single-item scales can reliably assess attitudes toward AI, including trust, risks, and benefits (Castro et al., 2023; Montag & Ali, 2025; Sindermann et al., 2021). For instance, Montag and Ali (2025) demonstrated the reliability of single-item measures for assessing trust in AI across cross-cultural settings, while Castro et al. (2023) found that single-item trust measures can be valid and reliable compared to multi-item scales. Similarly, the Attitudes Toward Artificial Intelligence (ATAI) scale includes a single item to measure trust as a component of positive attitudes (Sindermann et al., 2021), supporting the applicability of this approach for constructs like Potential Risks and Positive Changes. Future in-depth studies should employ multi-item, semantically rich scales with validated sub-dimensions, such as trust (cognitive vs. affective; Glikson & Woolley, 2020), ethical implications (idealist vs. pragmatic; Mökander & Floridi, 2023), and wellbeing (PERMA model; Butler & Kern, 2016).

Prior to responding, participants were presented with concrete, context-specific scenarios involving a cleaning robot and a smart car, each designed to vary in levels of automation and criticality. Within these scenarios, potential risks were embedded in the descriptions, for example, cleaning inefficiency in low-criticality contexts or safety-related consequences in high-criticality contexts, to provide a realistic basis for participants' evaluations. The Arabic version of these questions was developed using the recommended back-translation method (Brislin, 1970). Face validation trials with individuals from both populations were conducted for both the English and Arabic versions to ensure clarity and representativeness of the questions (Piçarra & Giger, 2018). Feedback indicated that certain scenarios needed a clearer depiction of autonomy levels. For instance, in the low versus high autonomy cleaning robot scenarios, the difference was not immediately apparent. To address this, we added a remote-control icon to the low autonomy scenario to visually emphasize the need for human intervention. Additionally, in the low autonomy scenario, we included a human figure to ensure it was clear that human control is required, even after automation. To make the criticality level indicator more comfortable for participants, we removed the word "risk" from the indicator picture. This adjustment aimed to focus participants' attention on the criticality aspect without inducing unnecessary discomfort. The four distinct scenarios, varying in automation and criticality levels as shown in Fig. 2, were presented to participants to illustrate the concepts. These scenarios included interactions with AI technologies like cleaning robots and smart cars, each described with varying degrees of automation (low and high) and criticality (low and high). The scenarios, as follows, were presented in a randomized order to avoid learning effect and order bias.

After each of the four scenarios, participants were asked to rate the

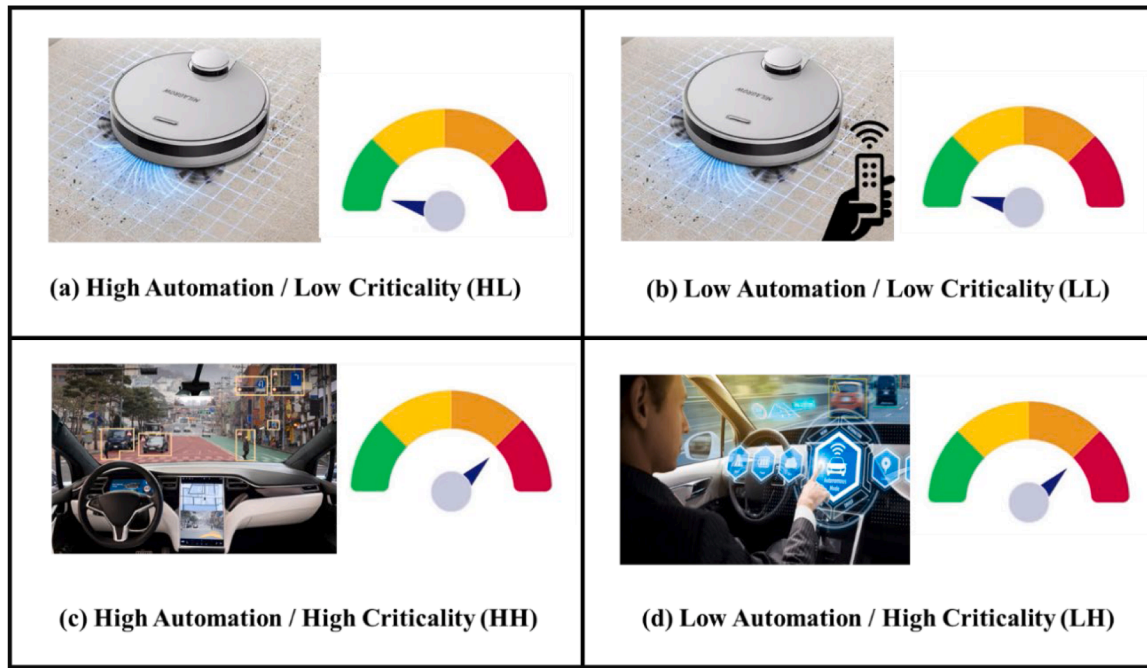


Fig. 2. Illustration of scenarios depicting various levels of automation and criticality in AI systems.

- a) **High Automation / Low Criticality (HL)**: AI operates independently in non-critical contexts.
- b) **Low Automation / Low Criticality (LL)**: Basic AI functions with significant human oversight
- c) **High Automation / High Criticality (HH)**: Nearly autonomous AI with significant potential impacts.
- d) **Low Automation / High Criticality (LH)**: AI decisions in critical situations require human guidance.

2.3. Data analysis

Descriptive analyses were calculated for UK and Arab samples. Skewness and kurtosis values were within the range of ± 2 for all variables except potential risk in scenario 4, which exhibited a kurtosis of 3.3. Pearson's correlation was used to assess the association between variables and Spearman's correlation was employed for the variable that violated the linearity assumption. No significant differences were observed between the two correlation methods. Hierarchical regression analysis was conducted to assess whether the independent variable, potential risk, the moderator's positive change, and their interaction could significantly predict users' perception of ethical implications. The assumptions of linearity and homoscedasticity were met. To assess the strength of the moderation, we used f^2 , to measure the proportion of variance in the dependent variable that is explained by the interaction term relative to the unexplained variance. The equation for f^2 is derived as follows:

$$f^2 = \frac{R_{included}^2 - R_{excluded}^2}{1 - R_{included}^2} \quad (1)$$

where, $R_{included}^2$ represents the R^2 value from the regression model that includes the interaction term (the full model) and $R_{excluded}^2$ represents the R^2 value from the regression model without the interaction term (the reduced model).

The effect sizes for each scenario were calculated following the method outlined by J Cohen (Kim, 2019). Furthermore, in multiple hypothesis testing scenarios, researchers often use the False Discovery Rate (FDR) to identify statistically significant results while controlling for the rate of false positives (Verhoeven et al., 2005). By effectively

controlling the rate of null hypotheses incorrectly rejected (Type I errors), FDR addresses the challenge posed by conducting numerous comparisons concurrently.

Before conducting the linear regression analysis, mean centering was applied to the predictor variables. Mean centering involves subtracting the mean of each variable from the corresponding variable's values, effectively transforming the data such that each variable has a mean of zero. This step is crucial for improving the interpretability of the regression coefficients and mitigating multicollinearity among the predictors. All statistical analyses were conducted using the JASP software package (J. Team, 2022).

3. Results

3.1. Descriptive statistics

Participants disclosed details regarding their age, gender, educational attainment, employment status, and country of residence. A tabulated overview of the demographic profiles of both Arab and UK cohorts is delineated in Table 1. As shown, the UK sample skewed slightly older and more female, while the Arab sample was younger, more gender-balanced, and had a higher proportion of participants with an undergraduate degree. Vocational/technical education was more common in the UK sample. Vocational/technical education is not as common in the Arab region compared to the UK. To provide a more accurate comparison, if we merge the percentages for vocational/technical education with those for a bachelor's degree (BSc) in the UK sample, the combined percentage would be approximately equal to the percentage of individuals with a BSc in the Arab sample.

Table 2 provides descriptive statistics for potential risks, positive changes, and ethical implications across various combinations of autonomy and criticality levels in both samples.

To minimize the influence of potential confounding from systematic differences in response styles between cultural groups, our main analyses focused on within-group comparisons rather than direct cross-group contrasts. Because we did not base our conclusions on between-group comparisons, any potential scale-use bias is unlikely to have substantially influenced the results. In addition, we examined the response distributions for each item in the Arab and UK samples (see

Table 1
Demographic characteristics of participants in the Arab and UK samples.

Variables		UK, N _{uk} = 316	Arab, N _{Arab} = 323
Gender	Male	96 (30.38 %)	147 (45.51 %)
	Female	220 (69.62 %)	176 (54.49 %)
Age	M (SD)	40.81 (10.35)	33.07 (9.05)
	Range	18 - 60	18 - 57
Education	No formal education	-	-
	Primary education (elementary)	2 (0.63 %)	2 (0.62 %)
	Secondary education (high school)	78 (24.68 %)	48 (14.86 %)
	Pursuing or completed vocational or technical education	71 (22.47 %)	13 (4.03 %)
	Pursuing or completed undergraduate degree (bachelor's)	104 (32.91 %)	221 (68.42 %)
Employment (%)	Pursuing or completed postgraduate degree (master's, Ph.D., etc.)	61 (19.30 %)	39 (12.07 %)
	Full-time employment	168 (53.16 %)	177 (54.80 %)
	Part-time employment	55 (17.40 %)	37 (11.45 %)
	Run my own business	15 (4.75 %)	21 (6.50 %)
	Homemaker	22 (6.96 %)	32 (9.91 %)
	Student	7 (2.22 %)	25 (7.74 %)
	Retired	11 (3.48 %)	7 (2.17 %)
	Unemployed	26 (8.23 %)	19 (5.88 %)
	Other	12 (3.80 %)	5 (1.55 %)

Figure S1 in Appendix). This inspection confirmed that both groups used the full range of the scale, with no consistent pattern of one group uniformly providing higher ratings across all measures.

3.2. Correlations between potential risks, positive changes, and ethical implications

The correlation analysis explores relationships between potential risks, positive changes, and ethical implications among participants from the UK and the Arab across various AI scenarios (LL, HL, LH, HH), which are shown in Table 3. In the Arab, a high positive correlation is observed between potential risks and ethical implications across all scenarios. In the HL scenario, this relationship is particularly high ($r = 0.66$, $p < 0.001$). Conversely, potential risks and positive changes

exhibit a significant negative correlation in all scenarios except for the LL scenario. Similar patterns were observed in the UK sample. For example, the positive correlation between potential risks and ethical implications in the HL scenario is $r = 0.64$ ($p < 0.001$), and the negative correlation between potential risks and positive changes in the same scenario is $r = -0.57$ ($p < 0.001$). Utilizing Pearson's provides a robust measure of linear relationships between these variables. Additionally, Spearman's correlations were calculated and included in the Appendix (Table S1) and there are no significant differences between the two correlations.

3.3. Multiple regression analysis

In the hierarchical linear regression analysis across four scenarios (LL, HL, LH, HH), potential risks consistently showed a strong and significant positive influence on the ethical implications of AI in the UK, with beta coefficients remaining robust across all models ($p < 0.001$) as shown in Table 4. Positive changes were introduced in the second step of the regression and generally exhibited a moderate negative impact on ethical implications, suggesting that positive perceptions can mitigate concerns (β values -0.26 for HH and -0.31 for LH, $p < 0.001$). Furthermore, interaction terms between potential risks and positive changes, added in the third step, indicated a moderating effect, slightly reducing the impact of potential risks on ethical concerns in LH and HH scenarios. This interaction was particularly notable in the high criticality scenarios (LH and HH), where positive change moderated the effect of potential

Table 3
Pearson's correlation analysis of potential risks, positive changes, and ethical implications among UK and Arab participants Across AI Scenarios.

Scenarios vs Questions	Variable	UK		Arab	
		Potential Risks	Positive Changes	Potential Risks	Positive Changes
LL	Positive Changes	-0.23***	-	-0.11	-
	Ethical Implications	0.67***	-0.10	0.55***	-0.18**
HL	Positive Changes	-0.56***	-	-0.24***	-
	Ethical Implications	0.64***	-0.37***	0.67***	-0.20***
LH	Positive Changes	-0.44***	-	-0.30***	-
	Ethical Implications	0.47***	-0.44***	0.53***	-0.29***
HH	Positive Changes	-0.48***	-	-0.33***	-
	Ethical Implications	0.46***	-0.42***	0.47***	-0.25***

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

Table 2
Descriptive statistics of perceptions of ai risks, positive changes, and ethical implications, among UK and Arab participants.

Variable			UK				Arab			
			M	SD	Skew	Kurt	M	SD	Skew	Kurt
Scenarios vs Questions	LL	Potential Risks	2.60	1.45	0.70	-0.44	3.42	1.58	-0.06	-1.30
		Positive Changes	4.17	1.22	-0.56	0.21	4.66	1.64	-0.92	1.25
		Ethical Implications	2.90	2.79	0.57	-0.73	4.83	3.07	-0.26	-1.04
	HL	Potential Risks	2.63	1.42	0.50	-0.81	3.35	1.65	-0.01	-1.38
		Positive Changes	4.66	1.11	-0.75	0.27	4.85	1.07	-0.76	0.73
		Ethical Implications	3.05	2.83	0.46	0.27	4.48	3.12	-0.06	-1.19
	LH	Potential Risks	4.44	1.25	-0.65	0.06	4.06	1.47	-0.59	-0.65
		Positive Changes	3.68	1.22	-0.37	-0.01	4.47	1.24	-0.83	0.39
		Ethical Implications	5.18	2.75	-0.16	-0.51	5.54	2.87	-0.45	-0.60
	HH	Potential Risks	5.16	1.08	-1.62	3.30	4.26	1.42	-0.63	-0.55
		Positive Changes	3.20	1.42	-0.01	-0.89	4.44	1.27	-0.74	0.08
		Ethical Implications	6.49	2.69	-0.64	-0.02	5.76	2.93	-0.49	-0.62

Note: Potential risks and Positive changes were rated on a scale from 1 ('Strongly disagree') to 6 ('Strongly agree'); Ethical implications were rated on a scale from 0 ('I'm not at all concerned') to 10 ('I'm completely concerned').

Table 4

Hierarchical linear regression analysis of potential risks, positive changes, and their interaction effects on ethical implications of AI in the UK sample.

			UK											
			Step 1				Step 2				Step 3			
Scenario	Variable		R ²	F	B	t	R ²	F	β	t	R ²	F	β	t
LL	Potential Risks		0.47	279.67***	0.69	16.72***	0.48	141.67***	0.70	16.66***	0.48	94.66***	0.69	15.53***
	Positive Changes								0.07	1.55			0.05	1.14
	Interaction of Potential Risks and Positive Changes												-0.04	-0.90
	Potential Risks		0.55	375.41***	0.74	19.38***	0.55	187.86***	0.72	15.58***	0.55	125.16***	0.72	15.50***
	Positive Changes								-0.04	-0.83			-0.04	-0.80
	Interaction of Potential Risks and Positive Changes												-0.03	-0.67
	Potential Risks		0.27	115.42***	0.52	10.74***	0.35	82.86***	0.39	7.61***	0.37	59.82***	0.41	8.06***
	Positive Changes								-0.31	-6.08***			-0.36	-6.82***
	Interaction of Potential Risks and Positive Changes												0.15	3.05**
	Potential Risks		0.30	131.47***	0.55	11.47***	0.35	82.60***	0.42	7.90***	0.36	57.52***	0.41	7.80***
	Positive Changes								-0.26	-4.89***			-0.32	-5.41***
	Interaction of Potential Risks and Positive Changes												0.12	2.26*

* p < 0.05

** p < 0.01

*** p < 0.001

risk on ethical implication, evidenced by a statistically significant increase in total variation explained of 9.7 % (R^2 change = 0.097), $F(3, 308) = 59.82$, $p < 0.001$ for LH and 6.2 % (R^2 change = 0.062), $F(3, 303) = 57.52$, $p < 0.001$ for HH, underscoring the complex dynamics between these factors and their implications for AI ethics and governance.

In the Arab sample, a hierarchical linear regression analysis across four AI scenarios (LL, HL, LH, HH) demonstrated that potential risks consistently had a strong and significant positive impact on the ethical implications of AI in the Arab. The beta coefficients remained robust across all models, with a significance level of $p < 0.001$, which is shown in Table 5. The introduction of positive changes in Step 2 revealed a generally negative influence on ethical implications, albeit with modest

effects in two scenarios ($\beta = -0.13$, $p < 0.01$ for LL; $\beta = -0.15$, $p < 0.01$) for LH). Notably, the interaction terms introduced in Step 3 highlighted a moderating effect of positive changes on the relationship between potential risks and ethical implications, especially in high-criticality scenarios (LH and HH), where the interaction terms were significant (R^2 change = 0.037, $F(3, 317) = 58.03$, $p < 0.01$ for LH; R^2 change = 0.042, $F(3, 318) = 36.41$, $p < 0.001$ for HH). These results underscore the importance of considering both potential risks and the potential mitigating effects of positive changes in shaping ethical evaluations of AI within the Arab context.

In the UK, in scenarios LH and HH, we observed small effect sizes

Table 5

Hierarchical linear regression analysis of potential risks, positive changes, and their interaction effects on ethical implications of AI in the Arab sample.

			Arab											
			Step 1				Step 2				Step 3			
Scenario	Variable		R ²	F	β	t	R ²	F	β	t	R ²	F	β	t
LL	Potential Risks		0.33	159.21***	0.58	12.62***	0.35	85.07***	0.56	12.39***	0.35	56.61***	0.56	12.28***
	Positive Changes								-0.13	-2.76**			-0.13	-2.77**
	Interaction of Potential Risks and Positive Changes												0.02	0.39
	Potential Risks		0.50	316.36***	0.71	17.79***	0.50	158.02***	0.70	17.06***	0.50	106.57***	0.69	16.61***
	Positive Changes								-0.02	-0.58			-0.03	-0.76
	Interaction of Potential Risks and Positive Changes												0.06	1.53
	Potential Risks		0.32	148.19***	0.56	12.17***	0.34	81.35***	0.52	10.87***	0.36	58.03***	0.49	10.30***
	Positive Changes								-0.15	-3.20**			-0.18	-3.73***
	Interaction of Potential Risks and Positive Changes												0.13	2.81**
	Potential Risks		0.21	87.75***	0.46	9.33***	0.22	46.00***	0.43	8.23***	0.26	36.41***	0.39	7.35***
	Positive Changes								-0.11	-2.03**			-0.16	-3.00**
	Interaction of Potential Risks and Positive Changes												0.19	3.69***

*p < 0.05

** p < 0.01

*** p < 0.001

($f^2 \approx 0.03$ and $f^2 \approx 0.02$ respectively), indicating some degree of interaction between the predictor variables and the moderator. Similarly, in the Arab, we observed similar patterns of effect sizes as scenarios LH: $f^2 \approx 0.03$ and HH: $f^2 \approx 0.03$ respectively), underscoring the importance of considering contextual factors in understanding moderation effects. In addition to observing small effect sizes in both the UK and Arab contexts, our findings align with the broader literature, where 72 % of reviewed studies demonstrated sufficient power to detect effects of 0.02, emphasizing the prevalence of similarly sized effects across various research endeavors (Aguinis et al., 2005).

Given these findings, in Table 6, we focus on LH and HH scenarios due to significant moderation effects observed in previous analyses. Effect sizes for moderating effects in these high-criticality scenarios offer valuable insights into the magnitude and practical significance of the interactions between potential risks and positive changes, which is shown in Table 6. In Table 6, potential risks consistently exerted a positive influence on ethical implications in both LH and HH scenarios for the UK. In the LH scenario, the 95 % confidence interval for β [0.93, 1.53] does not cross zero, indicating a significant positive effect ($\beta = 0.41$, $p < 0.001$). Similarly, in the HH scenario, the 95 % confidence interval for β [0.95, 1.60] confirms a positive influence ($\beta = 0.41$, $p < 0.001$). Conversely, potential risks consistently showed a positive influence on ethical implications across both LH and HH scenarios for Arab. For the LH scenario, the 95 % confidence interval for β [1.16, 1.70] remains above zero, demonstrating a significant positive effect ($\beta = 0.49$, $p < 0.001$). In the HH scenario, the 95 % confidence interval for β [0.83, 1.40] further supports this positive relationship ($\beta = 0.39$, $p < 0.001$). Moreover, the interaction between potential risks and positive changes demonstrated a significant moderating effect on ethical implications in both scenarios and regions, even after FDR correction.

4. Discussion

We examined the interplay between potential risks, positive changes, and ethical implications of AI across two culturally distinct samples. Consistent with theoretical expectations, some findings, such as heightened ethical concern in high-criticality scenarios, align with intuitive assumptions. However, in an era where AI has moved from institutional applications to everyday consumer use, with recent surveys showing that over half (53 %) of U.S. consumers now use or experiment with generative AI, and daily usage has become routine for many (Deloitte, 2025). Our study provides validation through a systematic, cross-cultural, scenario-based design. It reflects the types of AI systems that regular users encounter today, including autonomous cleaning robots and smart cars. By quantifying these effects across UK and Arab populations, we provide a foundation for future research to build upon with more complex and ecologically embedded designs. Our study reveals both similarities and nuanced differences across the two cultural contexts. The findings support Hypothesis H1, revealing a consistent positive relationship between potential risks and ethical implications of

AI across all tested scenarios (LL, HL, LH, HH) in both the UK and Arab regions.

Conversely, the negative impact of positive changes on ethical implications is evident from the substantial beta coefficients and significant p-values. In the UK's LH and HH scenarios, the significant relationship indicates a notable negative influence on ethical implications. Similarly, in the Arab region's LL and LH scenarios, the significant results reflect a meaningful negative effect. These findings suggest that positive changes tend to mitigate ethical concerns surrounding AI technologies. Therefore, Hypothesis H1 posits a negative relationship between positive changes and ethical implications.

Furthermore, the moderating role of positive changes, as outlined in Hypothesis H2, revealed context-specific effects. While the UK data showed some significant interactions, the Arab region demonstrated stronger and more consistent moderation effects, particularly in high-criticality scenarios. These findings suggest that positive changes may help alleviate ethical concerns in both contexts, potentially more so in the Arab region, where cultural factors could amplify sensitivity to risk (Pidgeon et al., 2008). Our results also indicate variability in risk perception, with Arab participants reporting higher potential risks in the LL and LH scenarios. This pattern may reflect a general sensitivity to uncertainty in lower-stakes, low-autonomy contexts, consistent with research showing that cultural dynamics significantly influence trust in automation (Chien et al., 2016). Conversely, lower risk ratings in HL and HH scenarios could suggest different interpretations of autonomy under high-stakes conditions, for example, viewing automated systems as more capable or less personally burdensome in critical situations. Our findings align with prior work emphasizing that user trust in AI is shaped by system-level features such as transparency, reliability, and perceived agency (Glikson & Woolley, 2020). These cultural dynamics warrant further investigation, as they may influence how individuals in different regions interpret risk, responsibility, and control in AI-mediated decision-making. Understanding these differences is essential for designing ethically aligned AI systems and communication strategies that resonate across cultural contexts (Stilgoe et al., 2013). We summarized the results in Table 7.

The positive relationship between potential risks and ethical implications of AI technologies corroborates findings from previous studies, which consistently highlight widespread concerns over AI's potential negative outcomes (Gînguță et al., 2023). These studies emphasize that as AI systems become more integrated into various aspects of daily life, public anxiety about their ethical implications increases. Extending beyond these established findings, our study elucidates those positive changes in AI perception, such as increased awareness of AI benefits or enhanced trust in self-governance, can significantly mitigate these concerns, particularly in high-risk scenarios (Roski et al., 2021). This mitigation effect was more consistently observed in scenarios involving high levels of automation and criticality, although the effect sizes were modest. Notably, our findings suggest that this moderating effect, where perceived positive changes reduce ethical concerns, may be more

Table 6
Effect Sizes for the Moderating Effects of Potential Risks and Positive Changes on Ethical Implications in High-Criticality Scenarios.

Scenario	Variable	UK							Arab						
		β	95 % CI for β		B	p	r^2_a	FDR	β	95 % CI for β		B	p	r^2_a	FDR
			Lower	Upper						Lower	Upper				
LH	Potential Risks	0.41	0.93	1.53	1.23	<0.001	0.14	0.02	0.49	1.16	1.70	1.43	<0.001	0.22	0.02
	Positive Changes	-0.36	-1.34	-0.74	-1.04	<0.001	0.10	0.03	-0.18	-0.79	-0.24	-0.52	<0.001	0.03	0.03
	Interaction of Potential Risks and Positive Changes	0.15	0.14	0.64	0.39	0.002	0.02	0.05	0.13	0.10	0.56	0.33	0.005	0.02	0.05
HH	Potential Risks	0.41	0.95	1.60	1.28	<0.001	0.13	0.02	0.39	0.83	1.43	1.01	<0.001	0.10	0.02
	Positive Changes	-0.32	-1.13	-0.53	-0.83	<0.001	0.06	0.03	-0.16	-0.77	-0.16	-0.42	0.003	0.02	0.03
	Interaction of Potential Risks and Positive Changes	0.12	0.04	0.59	0.32	0.024	0.01	0.05	0.19	0.24	0.80	0.46	<0.001	0.03	0.05

Table 7

Scenario-specific support for hypotheses regarding potential risks and positive changes in AI Ethics.

Hypothesis	Relationship	p-Value (UK)	Supported Scenarios (UK)	p-Value (Arab)	Supported Scenarios (Arab)
H1	Positive relationship between potential risks and ethical implications of AI	< 0.001	All scenarios (LL, HL, LH, HH)	< 0.001	All scenarios (LL, HL, LH, HH)
	Negative relationship between potential positive changes and ethical implications of AI.	Ranges from < 0.05 to < 0.001	Mostly all scenarios but weaker in some.	Ranges from < 0.05 to < 0.01	Less consistent, mostly weaker effects
H2	The moderating effect of potentially positive changes on the relationship between potential risks and ethical implications.	Ranges from non-significant to < 0.05	Weak effects, not consistently supported	Ranges from < 0.05 to < 0.001	More pronounced in LH, HH

pronounced in the Arab sample. This was particularly evident in high-criticality scenarios, where the presence of expected benefits more strongly attenuated the link between perceived risks and ethical implications. This regional variation may imply that cultural values and societal norms could influence how people weigh risks against benefits when evaluating AI technologies, as shown in comparative studies between Germany and China that reveal culturally distinct patterns in AI expectations, perceived benefits, and risk–benefit trade-offs (Brauner et al., 2024). It shows that individuals from individualist societies, such as Western countries, often perceive AI as a potential threat to personal autonomy, whereas collectivist cultures tend to integrate AI more harmoniously as a tool supporting social cohesion and conformity (Brauner et al., 2024).

These findings have potential implications for policymakers, industry professionals, and other stakeholders. However, because the study is exploratory and has methodological limitations, the results should be interpreted cautiously. The strong positive relationship between potential risks and ethical implications across all scenarios (LL, HL, LH, HH) in both the UK and Arab suggests that public apprehension regarding AI is nuanced. This relationship is stronger in scenarios where AI is viewed as critical, indicating that the degree of potential criticality of AI technologies amplifies ethical concerns. For policymakers, these findings may inform discussions about regulatory approaches that both address AI-related risks and acknowledge its potential benefits. For instance, the EU AI Act's risk-based framework, which classifies AI systems into categories from unacceptable to minimal risk (Future of Life Institute, 2024), shows how these regulatory discussions are developing. Our findings offer evidence for this tiered approach, especially the increased public concern we noted in high-criticality scenarios. This concern matches the Act's stricter requirements for 'high-risk' AI systems. Policies can aim to balance AI's potential for societal good, including advancements in healthcare diagnostics, environmental protection, and educational accessibility, as illustrated in scenarios where the positive impact was significant (Aggarwal et al., 2021; Hashiguchi et al., 2021). Industry professionals need to consider these perceptions, ensuring AI systems are transparent by marketing, and their benefits are communicated, an approach echoed in the findings from the highly positive change scenarios (Wind et al., 2019). Educational campaigns that elucidate successful AI use cases could play a crucial role in fostering informed and positive public opinion. This is particularly important in scenarios with neutral perceptions of positive change, where public opinion may be more malleable. To ensure these campaigns are effective, it is essential to strike a balance between explaining the risks associated with AI without inducing fear or withholding crucial details. This balance is similar to the challenge faced by healthcare professionals, such as doctors and nurses, who must deliver news to patients in a manner that is truthful yet not frightening or deceptive (R. F. Brown & Bylund, 2008). Research in healthcare communication strategies can provide valuable insights into achieving this balance in AI education campaigns (Albahri et al., 2018). Additionally, researchers are called upon for continued exploration of AI's perception across cultural contexts (Hofstede, 2001), considering the significant role of cultural factors in shaping perceptions as evidenced by the pronounced

effects in the LH and HH scenarios within the Arab.

Additionally, it is crucial to consider how the portrayal of AI's benefits can influence public acceptance and ethical evaluations. Behavioral economics emphasizes the importance of understanding human decision-making, particularly through concepts like mental accounting and other psychological factors (e.g., cognitive bias, emotional influences, etc.) that shape consumer behavior (Zhu, 2024). The study by Zhennan highlights the role of behavioral economics in shaping international entrepreneurial decisions, particularly through the integration of psychological factors such as cognitive biases and decision-making frameworks (Li et al., 2024). By understanding these psychological factors, we can better anticipate how individuals will react to new technologies such as AI. Public perception of AI is often influenced by how its benefits and risks are communicated, which can shape both acceptance and ethical evaluations. Framing the benefits of AI systems too positively may result in over-reliance while focusing solely on risks, which can lead to the rejection of valuable innovations. Therefore, a balanced approach is needed, where both the benefits and risks are presented transparently. In the context of marketing, ethics in AI play a crucial role, as misleading or exaggerated claims about AI's capabilities may not only erode trust but also hinder AI acceptance in the long term. Ethics in AI marketing further involves tackling critical issues such as privacy, data protection, algorithmic bias, and transparency. Addressing these challenges is fundamental to maintaining consumer trust and fostering responsible marketing practices, especially as AI increasingly plays a role in analyzing behavior, personalizing experiences, and automating decisions (Hari et al., 2025). While explainable AI (XAI) enhances transparency in AI systems yet poses risks of over-simplification that may lead to misunderstandings or mistrust among users, the study emphasizes the need for careful design and interdisciplinary collaboration between human-centred interaction (HCI) and AI (Reddy, 2024). To further support informed decision-making, AI introduction strategies should incorporate double-sided messages, presenting both advantages and drawbacks, providing a more comprehensive view of both the benefits and risks associated with AI technologies. The study by Meng supports this by showing that a double-sided message strategy enhances customers' willingness to engage with AI chatbots through the mediating role of perceived authenticity (Meng et al., 2023). Such balanced and ethical messaging could foster a more rational and measured approach to AI adoption, ensuring neither over-rejection nor over-reliance.

The use of a cross-sectional approach precludes the observation of changes over time, limiting our understanding of the evolution of perceptions and the directionality of the relationships examined. Future research could incorporate longitudinal or experimental designs to better assess causal pathways and examine how perceptions shift with repeated exposure or real-world AI interaction. Such approaches would also help uncover the moderating effects of demographic, psychological, or contextual variables over time. Additionally, a longitudinal approach and enhanced ecological validity are essential for capturing the complexity and lived experiences of participants, while acknowledging the need for a more nuanced exploration of cultural diversity within regions. Although participants were presented with detailed, context-

specific vignettes, our measure of “potential risks” relied on a single, broad item capturing perceived rather than theoretically delineated risk dimensions. Risk can be conceptualized and measured using established multidimensional frameworks, such as those distinguishing safety, financial, ethical, social, and health-related risks (Weber et al., 2002), or via subjective, participant-defined perceptions. Our study adopted the latter approach, assessing risk as perceived by respondents in the context of specific AI scenarios. This means our findings should be interpreted with caution, as comparisons with studies using domain-specific or objective risk measures may yield differing patterns. Our constructs of ‘potential risks’ and ‘positive changes’ may not capture all aspects of these multifaceted concepts. The complexity of ethical considerations may require more detailed and varied measures to fully encapsulate the range of public opinion.

The survey’s structure, which includes the wording of questions and the response options provided, may guide participants’ answers in specific directions due to the framing effect. The four distinct scenarios involving cleaning robots and smart cars were designed to illustrate concepts of automation and criticality. However, these specific scenarios might not fully represent the range of AI applications in real-world contexts, which could limit the generalizability of the findings to other AI technologies or settings. Although the four scenarios were carefully developed, face-validated, and systematically varied by automation and criticality levels, they represent a narrow subset of AI applications. This may limit the ecological validity and generalizability of the findings, as perceptions could differ substantially in domains that are more personally salient or socially consequential, such as healthcare, education, surveillance, or military contexts. For example, high-critical AI in healthcare may evoke stronger ethical concerns due to perceived risks to human life, while educational AI may elicit different evaluations linked to fairness or long-term societal impact.

Furthermore, while anchored vignettes and face validation helped reduce ambiguity and guided participant interpretation, individual differences in scenario comprehension may still exist. Some participants may have interpreted the automation levels differently despite visual cues, and this could impact the internal consistency of the measures.

A further limitation relates to the use of a non-probabilistic online panel (TGM Research). While this allowed efficient access to participants from the UK and Arab regions, it may introduce self-selection bias, as individuals who participate in online studies often differ from the broader population in terms of digital literacy, interest in AI, or motivation to engage in research (Bethlehem, 2010). Our study was designed to provide initial evidence of public attitudes toward AI across varying levels of automation and criticality, before expanding to larger and more resource-intensive studies with broader, stratified sampling and diverse recruitment channels (e.g., online and paper-based surveys).

5. Conclusion and future work

This study examined the relationship between perceived risks, potential positive changes, and ethical evaluations of AI across two cultural contexts, the UK and the Arab regions. The results indicate that higher perceived risks are consistently associated with greater ethical concerns across all tested scenarios (LL, HL, LH, HH), suggesting a consistent pattern irrespective of automation or criticality levels. Moreover, this association appears to be moderated by perceptions of positive change, which may help alleviate some ethical concerns, particularly in higher-risk scenarios where the influence of such perceptions was more evident.

Importantly, our findings show that cultural and regional context can influence how people weigh risks and benefits in forming ethical judgments, but that these effects vary by mode of operation. This pattern aligns with recent cross-cultural research by Brauner et al. (Brauner et al., 2024), who compared AI perceptions in Germany and China, finding that German participants emphasized risks over benefits, while Chinese respondents exhibited more balanced evaluations. Within both the UK and Arab samples, perceived risks consistently increased ethical

concern across all modes, while perceived positive changes mitigated these concerns more strongly in certain modes, particularly high-criticality scenarios in the Arab sample. Our results suggest that cultural context may influence how operational AI characteristics are interpreted, potentially in ways that align with Hofstede’s (Hofstede, 2001) dimensions such as uncertainty avoidance, collectivism, and power distance. This finding aligns with domain-specific risk research (Weber et al., 2002), demonstrating that risk tolerance varies systematically across both cultural contexts and technological domains, particularly when considering different levels of task stakes. These findings suggest that considering criticality–autonomy configurations can improve analyses of cross-cultural attitudes toward AI, and that governance and communication efforts may benefit from accounting for cultural context and the system’s operational mode.

Our analysis suggests that cultural and regional differences may influence perceptions of AI. While participants from both the UK and Arab regions acknowledged potential risks, their ethical evaluations varied, with Arab participants showing a somewhat stronger association with perceived positive changes. These patterns highlight the potential relevance of cultural context in informing ethical evaluations of AI and point to the value of developing AI policies and communication strategies that are responsive to cultural differences.

Future research could address several limitations to provide a more comprehensive understanding of societal attitudes toward AI technologies. It would benefit from including a broader demographic spectrum, encompassing individuals with varying levels of familiarity and engagement with AI. This approach would help capture a wider range of public opinions and provide a more nuanced view of societal attitudes towards AI. Additionally, employing longitudinal studies could offer insights into how perceptions and relationships evolve over time, addressing the limitations of the cross-sectional design used in the current study. Finally, examining the internal diversity within broad regions, such as the UK and Arab, would yield further insights into cultural nuances affecting ethical evaluations of AI. This granular cultural analysis could enhance the understanding of how different cultural contexts shape the perception of AI’s ethical implications. While participants reviewed context-specific scenarios involving cleaning robots and autonomous cars to help ground their judgments, the broad phrasing of “potential risks” remains a limitation. This approach captures participants’ subjective ethical evaluations in context but may also lead to variability based on how much the chosen domains resonate with their personal interests, experiences, or perceived relevance. Future research should refine this by distinguishing between specific risk domains (e.g., safety, financial, privacy) to yield more targeted insights, incorporating measures of domain familiarity or personal salience, and testing whether results differ across a broader set of AI applications beyond cleaning and autonomous driving. Such steps would help determine whether observed effects are scenario-specific or generalizable across technological domains.

Our findings underscore the importance of presenting a balanced view of AI to avoid over-reliance or rejection. Incorporating behavioral economics principles, particularly through the use of double-sided messages, can help ensure that users critically evaluate AI’s benefits and risks. Future AI designs may benefit from looking at these frameworks to encourage responsible technology and sustainable use in various cultural contexts.

Data availability statements

The study design, the dataset, and the data dictionary can be found at the following Open Science Framework link: <https://osf.io/rmqqt2/>

Ethical conduct

The study was conducted in accordance with the ethical standards set forth in the Declaration of Helsinki and was approved by the

Institutional Review Board (IRB) of the first author (ID: IRB-2024-59).

Declaration of generative AI and AI-assisted technologies in the writing process

During the preparation of this work the authors used ChatGPT in order correct language errors and suggest restyling when the language was deemed difficult to read by some authors. After using this tool/service, the authors reviewed and edited the content as needed and take full responsibility for the content of the publication.

Funding

This publication was supported by NPRP 14 Cluster grant number NPRP 14C-0916-210015 from the Qatar National Research Fund (a member of Qatar Foundation). The findings herein reflect the work and are solely the responsibility of the authors.

CRediT authorship contribution statement

Mohammad Mominur Rahman: Writing – original draft, Methodology, Investigation, Formal analysis, Conceptualization. **Areej Babiker:** Writing – review & editing, Validation, Methodology, Investigation, Data curation, Conceptualization. **Ala Yankouskaya:** Writing – review & editing, Validation, Supervision, Methodology, Investigation, Formal analysis, Conceptualization. **Sameha AlShakhsi:** Writing – review & editing, Data curation. **Raian Ali:** Writing – review & editing, Supervision, Project administration, Methodology, Funding acquisition, Data curation, Conceptualization.

Declaration of competing interest

The authors declare no conflict of interest in this work.

Supplementary materials

Supplementary material associated with this article can be found, in the online version, at [doi:10.1016/j.jrt.2026.100163](https://doi.org/10.1016/j.jrt.2026.100163).

References

- Aggarwal, M., Gingras, C., & Deber, R. (2021). Artificial intelligence in healthcare from a policy perspective. *Multiple Perspectives on Artificial Intelligence in Healthcare: Opportunities and Challenges*, 53–64. https://doi.org/10.1007/978-3-030-67303-1_5
- Aguinis, H., Beaty, J. C., Boik, R. J., & Pierce, C. A. (2005). Effect size and power in assessing moderating effects of categorical variables using multiple regression: a 30-year review. *Journal of Applied Psychology*, 90(1), 94. <https://doi.org/10.1037/0021-9010.90.1.94>
- Ajenaghughrur, I. Ben, da Costa Sousa, S. C., & Lamas, D. (2020). Risk and trust in artificial intelligence technologies: A case study of autonomous vehicles. In *2020 13th International Conference on Human System Interaction (HSI)* (pp. 118–123). <https://doi.org/10.1109/HSI49210.2020.9142686>
- Albahri, A. H., Abushibs, A. S., & Abushibs, N. S. (2018). Barriers to effective communication between family physicians and patients in walk-in centre setting in Dubai: A cross-sectional survey. *BMC Health Services Research*, 18(637), 1–13. <https://doi.org/10.1186/s12913-018-3457-3>
- Alexander, C. S., & Becker, H. J. (1978). The use of vignettes in survey research. *Public Opinion Quarterly*, 42(1), 93–104.
- Atzmüller, C., & Steiner, P. M. (2010). Experimental vignette studies in survey research. *Methodology*.
- Awad, E., Desouza, S., Kim, R., Schulz, J., Henrich, J., Shariff, A., Bonnefon, J.-F., & Rahwan, I. (2018a). The moral machine experiment. *Nature*, 563(7729), 59–64.
- Awad, E., Desouza, S., Kim, R., Schulz, J., Henrich, J., Shariff, A., Bonnefon, J.-F., & Rahwan, I. (2018b). The moral machine experiment. *Nature*, 563(7729), 59–64. <https://doi.org/10.1038/s41586-018-0637-6>
- Azuma, M. C., Giordano, F. B., Stoffregen, S. A., Klos, L. S., & Lee, J. (2023). It practically drives itself: autonomous vehicle technology, psychological attitudes, and susceptibility to risky driving behaviors. *Ergonomics*, 66(2), 246–260. <https://doi.org/10.1080/00140139.2022.2076906>
- Baek, C., & Doleck, T. (2024). Learning analytics: A comparison of Western, educated, industrialized, rich, and Democratic (WEIRD) and Non-WEIRD research. *Knowledge Management & E-Learning*, 16(2), 217–236.
- Banks, V., Shaw, E., & Large, D. R. (2019). Keeping the driver in the loop: the 'other' ethics of automation. In *Proceedings of the 20th Congress of the International Ergonomics Association (IEA 2018) Volume VI: Transport Ergonomics and Human Factors (TEHF)*, Aerospace Human Factors and Ergonomics 20 (pp. 70–79). https://doi.org/10.1007/978-3-319-96074-6_8
- Bethlehem, J. (2010). Selection bias in web surveys. *International Statistical Review*, 78(2), 161–188. <https://doi.org/10.1111/j.1751-5823.2010.00112.x>
- Bigman, Y. E., & Gray, K. (2018). People are averse to machines making moral decisions. *Cognition*, 181, 21–34.
- Brauner, P., Glawe, F., Liehner, G. L., Vervier, L., & Ziefle, M. (2024). AI perceptions across cultures: Similarities and differences in expectations, risks, benefits, tradeoffs, and value in Germany and China. *ArXiv Preprint ArXiv:2412.13841*.
- Brislin, R. W. (1970). Back-translation for cross-cultural research. *Journal of Cross-Cultural Psychology*, 1(3), 185–216. <https://doi.org/10.1177/135910457000100301>
- Brossard, D., & Scheufele, D. A. (2013). Science, new media, and the public. *Science (New York, N.Y.)*, 339(6115), 40–41. <https://doi.org/10.1126/science.1232329>
- Brown, R. F., & Bylund, C. L. (2008). Communication skills training: describing a new conceptual model. *Academic Medicine*, 83(1), 37–44. <https://doi.org/10.1097/ACM.0b013e31815c631e>
- Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., & Askell, A. (2020). Language models are few-shot learners. *Advances in Neural Information Processing Systems*, 33, 1877–1901. <https://doi.org/10.48550/arXiv.2005.14165>
- Bruin, J. (2011). Newtest: command to compute new test. <https://stats.oarc.ucla.edu/sas/modules/introduction-to-the-features-of-sas/>
- Butler, J., & Kern, M. L. (2016). The PERMA-Profil: A brief multidimensional measure of flourishing. *International journal of wellbeing*, 6(3).
- Calvaresi, D., Carli, R., Tiribelli, S., Buzcu, B., Aydogan, R., Di Vincenzo, A., Mualla, Y., Schumacher, M., & Calbimonte, J. P. (2025). Computational persuasion technologies, explainability, and ethical-legal implications: A systematic literature review. In *Computers in human behavior reports*, 17. Elsevier B.V. <https://doi.org/10.1016/j.chbr.2024.100577>
- Camerer, C. F. (2019). Artificial intelligence and behavioral economics. *The Economics of Artificial Intelligence: An Agenda*, 587–608. <https://doi.org/10.7208/9780226613475-026>
- Castro, et al. (2023). Comparing single-item and multi-item trust scales: Insights for assessing trust in project leaders. *Behavioral Sciences*, 13(9), 786. <https://doi.org/10.3390/BS13090786>. Vol. 13, Page 786 Sep. 2023.
- Chanseau, A., Dautenhahn, K., Walters, M., Lakatos, G., Koay, K. L., & Salem, M. (2017). People's perceptions of task criticality and preferences for robot autonomy. *Poster Papers*, 65. <https://doi.org/10.31256/ukras17.21>
- Chien, S.-Y., Lewis, M., Sycara, K., Liu, J.-S., & Kumru, A. (2016). Influence of cultural factors in dynamic trust in automation. In *2016 IEEE International Conference on Systems, Man, and Cybernetics (SMC)* (pp. 2884–2889).
- Choudhury, A. (2022). Toward an ecologically valid conceptual framework for the use of artificial intelligence in clinical settings: Need for systems thinking, accountability, decision-making, trust, and patient safety considerations in safeguarding the technology and clinicians. *JMIR Hum Factors*, 9(2), Article e35421. <https://doi.org/10.2196/35421>
- Cushman, F. (2013). Action, outcome, and value: A dual-system framework for morality. *Personality and Social Psychology Review*, 17(3), 273–292.
- Davenport, T. H., & Ronanki, R. (2018). Artificial intelligence for the real world. *Harvard Business Review*, 96(1), 108–116.
- Deloitte. (2025). Deloitte: As generative AI gains ground, consumers choose the innovators they trust. <https://www.deloitte.com/us/en/about/press-room/connectivity-mobile-trends-survey.html>
- Dignum, V. (2017). Responsible artificial intelligence: designing AI for human values. *ITU Journal: ICT Discoveries*, 1.
- Elgammal, A., Liu, B., Elhoseiny, M., & Mazzone, M. (2017). Can: creative adversarial networks, generating "art" by learning about styles and deviating from style norms. *ArXiv Preprint ArXiv:1706.07068*, 6. <https://doi.org/10.48550/arXiv.1706.07068>
- Figueras, C., Verhagen, H., & Cerratto Pargman, T. (2021). Trustworthy AI for the people? In *Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society* (pp. 269–270). <https://doi.org/10.1145/3461702.3462470>
- Future of Life Institute. (2024). High-level summary of the ai act. February 27. artificialintelligenceact.eu/high-level-summary/. Retrieved March 2, 2026, from <https://artificialintelligenceact.eu/high-level-summary/>
- Gînguță, A., Ștefău, P., Noja, G. G., & Munteanu, V. P. (2023). Ethical impacts, risks and challenges of artificial intelligence technologies in business consulting: A new modelling approach based on structural equations. *Electronics*, 12(6), 1462. <https://doi.org/10.3390/electronics12061462>
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660. <https://doi.org/10.5465/annals.2018.0057>
- Greene, J. D., Somerville, R. B., Nystrom, L. E., Darley, J. M., & Cohen, J. D. (2001). An fMRI investigation of emotional engagement in moral judgment. *Science (New York, N.Y.)*, 293(5537), 2105–2108.
- Greg Laski. (n.d.). MOBILE panel sample and online surveys tgm research. Accessed: Jan. 09, 2024. [Online]. Available: <https://Tgmresearch.Com/>
- Hari, H., Sharma, A., Verma, S., & Chaturvedi, R. (2025). Exploring ethical frontiers of artificial intelligence in marketing. *Journal of Responsible Technology*, 21. <https://doi.org/10.1016/j.jrt.2024.100103>
- Hashiguchi, T. C. O., Slawomirski, L., & Oederkirk, J. (2021). Laying the foundations for artificial intelligence in health. *OECD Health Working Papers, OECD Publishing*, 128. <https://doi.org/10.1787/3f62817d-en>

- Hofstede, G. (2001). *Culture's consequences: comparing values, behaviors, institutions and organizations across nations*, 41. Sage publications. [https://doi.org/10.1016/S0005-7967\(02\)00184-5](https://doi.org/10.1016/S0005-7967(02)00184-5)
- Hrnjic, E., & Tomczak, N. (2019). Machine learning and behavioral economics for personalized choice architecture. *ArXiv Preprint ArXiv:1907.02100*. <https://doi.org/10.48550/arXiv.1907.02100>
- Huriye, A. Z. (2023). The ethics of artificial intelligence: examining the ethical considerations surrounding the development and use of AI. *American Journal of Technology*, 2(1), 37–44. <https://doi.org/10.58425/ajt.v2i1.142>
- J. Team. (2022). *JASP - A fresh way to do statistics*. JASP [Online]. Available <https://Jasp-Stats.Org/>.
- Kamila, M. K., & Jasrotia, S. S. (2023). Ethical issues in the development of artificial intelligence: Recognizing the risks. *International Journal of Ethics and Systems*. <https://doi.org/10.1108/IJOES-05-2023-0107>
- Kamilaris, A., & Prenafeta-Boldú, F. X. (2018). Deep learning in agriculture: A survey. *Computers and Electronics in Agriculture*, 147, 70–90.
- Kaswan, K. S., Dhatteerwal, J. S., & Ojha, R. P. (2024). AI in personalized learning. *Advances in Technological Innovations in Higher Education*, 103–117. <https://doi.org/10.1201/9781003376699-9>
- Kieslich, K., Keller, B., & Starke, C. (2022). Artificial intelligence ethics by design. Evaluating public perception on the importance of ethical design principles of artificial intelligence. *Big Data & Society*, 9(1), Article 20539517221092956. <https://doi.org/10.1177/20539517221092956>
- Kim, J. H. (2019). Multicollinearity and misleading statistical results. *Korean Journal of Anesthesiology*, 72(6), 558. <https://doi.org/10.4097/kja.19087>
- Klein, U., Depping, J., Wohlfahrt, L., & Fassbender, P. (2024). Application of artificial intelligence: Risk perception and trust in the work context with different impact levels and task types. *AI & SOCIETY*, 39(5), 2445–2456. <https://doi.org/10.1007/s00146-023-01699-w>
- Li, Z., Shu, K., Cheng, M., & Nian, J. (2024). Exploring the nexus between international perspective, innovation resources, and entrepreneurial internationalization: A resource-based and behavioral perspective. *Journal of the Knowledge Economy*, 1–34. <https://doi.org/10.1007/s13132-024-02012-w>
- Linzen, S., Sturm, C., Brühlmann, F., Cassau, V., Opwis, K., & Reinecke, K. (2021). How weird is CHI? In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems* (pp. 1–14).
- Ma, Y., Wang, Z., Yang, H., & Yang, L. (2020). Artificial intelligence applications in the development of autonomous vehicles: A survey. *IEEE/CAA Journal of Automatica Sinica*, 7(2), 315–329.
- Meng, L. M., Li, T., & Huang, X. (2023). Double-sided messages improve the acceptance of chatbots. *Annals of Tourism Research*, 102(1), Article 103644. <https://doi.org/10.1016/j.annals.2023.103644>
- Mökander, J., & Floridi, L. (2023). Operationalising AI governance through ethics-based auditing: an industry case study. *AI and ethics*, 3(2), 451–468. <https://doi.org/10.1007/s43681-022-00171-7>
- Mollen, J. (2024). Towards a research ethics of real-world experimentation with emerging technology. *Journal of Responsible Technology*, 20, Article 100098. <https://doi.org/10.1016/j.jrt.2024.100098>
- Montag, C., Nakov, P., & Ali, R. (2024). Considering the IMPACT framework to understand the AI-well-being-complex from an interdisciplinary perspective. *Telematics and Informatics Reports*, 13, Article 100112.
- Montag, C., & Ali, R. (2025). Can we assess attitudes toward AI with single items? Associations with existing attitudes toward AI measures and trust in ChatGPT. *J. technol. behav. sci.*. <https://doi.org/10.1007/s41347-025-00481-7>
- Morley, J., Machado, C. C. V., Burr, C., Cows, J., Joshi, I., Taddeo, M., & Floridi, L. (2020). The ethics of AI in health care: A mapping review. *Social Science & Medicine*, 260, Article 113172. <https://doi.org/10.1016/j.socscimed.2020.113172>
- Mulgan, T. (2016). Superintelligence: paths, dangers, strategies. In *The philosophical quarterly*, 66. Oxford University Press. <https://doi.org/10.1093/pq/pqv034>
- Ouchchy, L., Coin, A., & Dubljević, V. (2020). AI in the headlines: The portrayal of the ethical issues of artificial intelligence in the media. *AI & SOCIETY*, 35, 927–936. <https://doi.org/10.1007/s00146-020-00965-5>
- Parasuraman, R., Sheridan, T. B., & Wickens, C. D. (2000). A model for types and levels of human interaction with automation. *IEEE Transactions on Systems, Man, and Cybernetics-Part A: Systems and Humans*, 30(3), 286–297.
- Park, J., & Jung, Y. (2021). Employees' intention to use artificial intelligence: Roles of perceived usefulness, trust, and perceived organizational support. *Korean Journal of Industrial and Organizational Psychology*, 34(2), 183–211. <https://doi.org/10.24230/kjiop.v34i2.183-211>
- Parycek, P., Schmid, V., & Novak, A.-S. (2024). Artificial intelligence (AI) and automation in administrative procedures: Potentials, limitations, and framework conditions. *Journal of the Knowledge Economy*, 15(2), 8390–8415. <https://doi.org/10.1007/s13132-023-01433-3>
- Piçarra, N., & Giger, J.-C. (2018). Predicting intention to work with social robots at anticipation stage: Assessing the role of behavioral desire and anticipated emotions. *Computers in Human Behavior*, 86, 129–146. <https://doi.org/10.1016/j.chb.2018.04.026>
- Pidgeon, N., Simmons, P., Sarre, S., Henwood, K., & Smith, N. (2008). *The ethics of socio-cultural risk research*, 10 pp. 321–329. <https://doi.org/10.1080/13698570802334526>
- Rahman, M. M., Babiker, A., & Ali, R. (2024). Motivation, concerns, and attitudes towards AI: Differences by gender, age, and culture. In *International Conference on Web Information Systems Engineering* (pp. 375–391).
- Reddy, S. T. A. (2024). Human-computer interaction techniques for explainable artificial intelligence systems. *Recent Trends in Artificial Intelligence & Its Applications*, 3(1), 1–7. <https://doi.org/10.46610/RTAIA.2024.v03i01.001>
- Roski, J., Maier, E. J., Vigilante, K., Kane, E. A., & Matheny, M. E. (2021). Enhancing trust in AI through industry self-governance. *Journal of the American Medical Informatics Association*, 28(7), 1582–1590. <https://doi.org/10.1093/jamia/ocab065>
- Septiandri, A. A., Constantinides, M., Tahaei, M., & Quercia, D. (2023). WEIRD FAcTs: How western, educated, industrialized, rich, and democratic is FAcT? In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency* (pp. 160–171).
- Sindermann, C., Sha, P., Zhou, M., et al. (2021). Assessing the attitude towards artificial intelligence: introduction of a short measure in German, Chinese, and english language. *Künstl Intell*, 35, 109–118. <https://doi.org/10.1007/s13218-020-00689-0>
- Stahl, B. C. (2021). *Artificial intelligence for a better future: an ecosystem perspective on the ethics of ai and emerging digital technologies*. Springer Nature. <https://doi.org/10.1007/978-3-030-69978-9>
- Steiner, P. M., Atzmüller, C., & Su, D. (2016). Designing valid and reliable vignette experiments for survey research: A case study on the fair gender income gap. *Journal of Methods and Measurement in the Social Sciences*, 7(2), 52–94.
- Stilgoe, J., Owen, R., & Macnaghten, P. (2013). Developing a framework for responsible innovation. *Research Policy*, 42(9), 1568–1580. <https://doi.org/10.1016/j.respol.2013.05.008>
- Suleyman, M. (2023). *The coming wave: technology, power, and the twenty-first century's greatest dilemma*, 1. Crown. <https://doi.org/10.1080/15228053.2023.2286781>
- Tolmeijer, S., Christen, M., Kandul, S., Kneer, M., & Bernstein, A. (2022). Capable but amoral? Comparing AI and human expert collaboration in ethical decision making. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems* (pp. 1–17). <https://doi.org/10.1145/3491102.3517732>
- Verhoeven, K. J. F., Simonsen, K. L., & McIntyre, L. M. (2005). Implementing false discovery rate control: Increasing your power. *Oikos (Copenhagen, Denmark)*, 108(3), 643–647. <https://doi.org/10.1111/j.0030-1299.2005.13727.x>
- von Garrel, J., & Jahn, C. (2023). Design framework for the implementation of AI-based (Service) business models for small and medium-sized manufacturing enterprises. *Journal of the Knowledge Economy*, 14(3), 3551–3569. <https://doi.org/10.1007/s13132-022-01003-z>
- Waclawski, E. (2012). How I use it: survey monkey. *Occupational Medicine*, 62(6), 477. <https://doi.org/10.1093/occmed/kqs075>
- Wanner, J., Herm, L.-V., Heinrich, K., & Janiesch, C. (2022). The effect of transparency and trust on intelligent system acceptance: Evidence from a user-based study. *Electronic Markets*, 32(4), 2079–2102. <https://doi.org/10.1007/s12525-022-00593-5>
- Ward, C., Raue, M., Lee, C., D'Ambrosio, L., & Coughlin, J. F. (2017). Acceptance of automated driving across generations: The role of risk and benefit perception, knowledge, and trust. *Human-computer interaction*. In , 19. *User Interface Design, Development and Multimodality: 19th International Conference, HCI International 2017, Vancouver, BC, Canada, July 9-14, 2017, Proceedings, Part I* (pp. 254–266). https://doi.org/10.1007/978-3-319-58071-5_20
- Weber, E. U., Blais, A., & Betz, N. E. (2002). A domain-specific risk-attitude scale: Measuring risk perceptions and risk behaviors. *Journal of Behavioral Decision Making*, 15(4), 263–290.
- Wilson, T. D., Lindsey, S., & Schooler, T. Y. (2000). A model of dual attitudes. *Psychological Review*, 107(1), 101.
- Wind, A., Constantinides, E., & de Vries, S. (2019). Marketing a transparent artificial intelligence (AI): A preliminary study on message design. In *18th International Marketing Trends Conference 2019*.
- Wu, H., Liu, J., & Liang, B. (2024). AI-driven supply chain transformation in industry 5.0: Enhancing resilience and sustainability. *Journal of the Knowledge Economy*. <https://doi.org/10.1007/s13132-024-01999-6>
- Xiong, Y., Shi, Y., Pu, Q., & Liu, N. (2023). More trust or more risk? User acceptance of artificial intelligence virtual assistant. *Human Factors and Ergonomics in Manufacturing & Service Industries*, 34, 190–205. <https://doi.org/10.1002/hfm.21020>
- Zhu, J. (2024). *Behavioral economics*, 1. Dean & Francis Press.