

Jakub Fryderyk Juranek , Grzegorz Pochwatko *Institute of Psychology, Polish Academy of Sciences, Warsaw, Poland*

Corresponding author – Jakub Fryderyk Juranek – jakub.juranek@sd.psych.pan.pl

Measurement of General Attitude Toward AI and its Relationship with AI Anxiety, Attitude Toward Robots and Belief in Human Nature Uniqueness: Polish Adaptation of the ATTARI-12 Questionnaire.

Abstract: Objective: This study introduces the ATTARI-12-PL, a Polish adaptation of the ATTARI-12 questionnaire designed to measure general attitudes toward artificial intelligence (AI) across cognitive, affective, and behavioral facets. Method: An online sample of 332 participants was used for the psychometric evaluation of the ATTARI-12-PL scale. The internal structure and consistency of the scale were assessed, and the validity of the questionnaire was evaluated through correlation analyses. Results: Single-factor model did not provide adequate fit in confirmatory factor analysis (CFA). In contrast, a bifactor model incorporating a general attitude factor alongside specific orthogonal factors (attitude facets and item wording) demonstrated very good fit indices and accounted well for the data structure. The scale showed high internal consistency (Cronbach's $\alpha = 0.94$). Strong negative correlations were found between AI attitude and negative attitude toward interacting with robots, moderate negative correlations with AI anxiety and negative attitude toward robots with human traits, and a mild negative association with belief in human nature uniqueness (BHNU). The scale also showed a strong positive correlation with other AI measurement tool: AIAS-4-PL. Conclusions: The ATTARI-12-PL is a reliable and valid tool for measuring general attitudes toward AI in the Polish language context. The scale's properties and significant correlations with related constructs support its use in future research on AI attitudes.

Keywords: *artificial intelligence, attitude toward AI, attitude toward robots, AI anxiety, belief in human nature uniqueness, humanoid robots*

Since the inception of the term ‘Artificial Intelligence’ (AI) during the Dartmouth Summer Research Project on Artificial Intelligence in 1956, led by John McCarthy (Chudziak et al., 2022; Kubassova et al., 2021), the associated technology has undergone several decades of both rapid advancement and occasional slowdowns (referred to as ‘AI winters’). These phases have culminated in remarkable achievements and widespread adoption of AI-powered technologies experienced today.

Artificial Intelligence is increasingly integrated into various aspects of modern societies, including education (Chen et al., 2020; Juranek, 2026), transportation

(Bharadiya, 2023), finance (Jáuregui-Velarde et al., 2024), farming (Holzinger et al., 2024), science (Xu et al., 2021), healthcare (Secinaro et al., 2021) and assistive technology (Juranek, 2025; Schaur & Matausch-Mahr, 2025). AI technology brings benefits and risks, of which some have already materialized and other are projected to be encountered in the future. Currently broadly recognized challenges considering AI include algorithmic bias (Kordzadeh & Ghasemaghaei, 2022), safety and malicious use worries (Blauth et al., 2022), copyright issues (Lee et al., 2024) and global job loss concerns (Khogali & Mekid, 2023).



The actual advances and risks associated with AI systems development had been for long time preceded by forecasts and projections present broadly not only in scientific literature but in popular culture, especially in the science-fiction genre. Cave and Dihal (2019; Watts & Bode, 2024) analyzed 300 works of fiction and non-fiction and distinguished four main categories of oppositions (dichotomies) expressed in imaginings of AI, each with a positive pole, called a 'hope' and its opposite, called 'fear'. The first dichotomy is the hope of extended life for humans (hope of 'immortality') accompanied by the fear of losing human identity ('inhumanity'). Another pair of opposites is the hope of life free of work ('ease') with contrasting fear of becoming redundant due to emergence of very capable AI ('obsolescence'). The third hope is rooted in the expectation that AI can extraordinarily fulfil individual's desires ('gratification'), accompanied by the fear of alienation: humans becoming redundant to each other. The last dichotomy is based on the hope of AI offering great power over others ('dominance'), alongside the fear that it will turn against humanity ('uprising'). The authors argue, that regardless of how these recognized, common perceptions of AI possibilities relate to the actual realities of the technology is of less importance, than how these perspectives can actually strongly influence how the real-world AI systems are seen, developed, deployed and regulated.

The theoretical work on perceptions of AI is accompanied by empirical approaches aimed of conceptualizing and measuring the actual attitudes toward AI among people. Attitudes comprise general assessments of individuals, groups, topics and objects. They play a role in the formation of beliefs (Ajzen, 2001) and shape decisions and behaviors (Glasman & Albarracín, 2006). Attitudes vary in terms of their origin, having affective, cognitive or behavioral grounding (Fishbein & Ajzen, 1977; Ostrom, 1969; Svenningsson et al., 2022), they also vary in the degree of strength - certain attitudes have greater impact and better predict behavior than others. In the assessment of attitude strength, consideration is given to indicators evaluated as more objective or subjective. Objectively, attitude strength is reflected in stability (remaining consistent over time), resistance (withstanding challenges), accessibility (quick recall), and spread (influencing related thoughts and behaviors). Subjective indicators include attitude certainty and its perceived moral basis (Luttrell & Sawicki, 2020). Attitudes are sometimes being assessed as implicit, automatic and unconscious evaluations, but the predominant methods of their measurement rely on people's responses to self-report measures (i.e. questionnaires), which approach is known as explicit attitude measurement (Bohner & Dickel, 2011; Briñol et al., 2019).

Several measures have been developed so far to measure attitudes toward AI. These are the General Attitudes Towards Artificial Intelligence Scale (GAAIS; Schepman & Rodway, 2020), the Attitudes Towards Artificial Intelligence Scale (ATAI; Sindermann et al., 2021), the AI Attitude Scale (AIAS-4; Grassini, 2023) and the ATTARI-12 questionnaire (Stein et al., 2024). These

instruments were created to evaluate AI attitude in general. Furthermore, there are scales specifically designed to measure anxiety, concerns, and perceived threats associated with AI or related technology, such as robots and autonomous agents. These comprise the Concerns About Autonomous Technology questionnaire (Stein et al., 2019), the Negative Attitude Toward Robots Scale (NARS; Nomura et al., 2006), the Threats of Artificial Intelligence Scale (TAI; Kieslich et al., 2021), and the AI Anxiety Scale (AIAS; Wang & Wang, 2022).

One of the main studied questions regarding attitudes toward AI is its determinants. The identified factors vary and include personality traits (Park & Woo, 2022), demographic variables and sociocultural dimensions (Méndez-Suárez et al., 2024; Rice & Wu, 2025), like media use habits (Brewer et al., 2022), level of exposure to technology or knowledge about it (Kaya et al., 2022). Research on personality predictors of attitudes toward AI have been instructive yet sometimes conflicting. Stein and colleagues (2024) reported that among the Big Five traits, Agreeableness consistently predicted more positive views toward AI. Similar findings for Agreeableness were observed by Guingrich and Graziano (2024), as well as by Sindermann and colleagues in a Chinese sample (Sindermann et al., 2021), where additionally Openness emerged as a significant positive predictor of AI acceptance. In contrast, Kaya et al. (2022) found that in a Turkish sample, Agreeableness was linked to more negative attitudes toward AI, highlighting the potential influence of cultural context. Schepman and Rodway (2023) in their studies demonstrated that Introversion, Conscientiousness, and Agreeableness were linked to more favorable AI attitudes, particularly in relation to a specific notion of forgiving negative aspects of AI. Neuroticism has repeatedly been associated with more negative outlook toward AI technologies (Guingrich & Graziano, 2024; Stănescu & Romaşcanu, 2024).

Park and Woo (2022) provided a more complex and nuanced perspective by examining how personality traits along with personal innovativeness in information technology (PIIT) predict attitudes toward AI across two affective (positive and negative emotions) and cognitive (sociality and functionality) dimensions. Their study exhibited that prior use and experience with AI serves as proximal predictor of attitudes, showing that both stable personality traits and experiential factors jointly shape how individuals evaluate AI. Taken together, the available body of work suggests that Agreeableness and Openness are most often traits that foster positive AI attitudes, Neuroticism tends to predict negativity, and contextual or experiential variables such as cultural background and prior exposure to AI significantly moderate these relationships.

In the Polish language context, a very limited psychometrical tooling associated with measuring attitudes toward AI was available at the time here reported project was initiated. In 2023, Fortuna and colleagues created the adaptation of the AIAS AI anxiety scale (AIAS-PL, Fortuna et al., 2023; Modliński et al., 2023), earlier the Negative Attitude Toward Robots Scale has been adapted

by Pochwatko and colleagues (NARS-PL; Pochwatko et al., 2015). In 2023, a cross-national study involving a Polish sample reported the utilization of a shortened GAAIS (Schepman & Rodway, 2023) questionnaire in the Polish language for the purpose of comprehensive research; however, the survey was originally designed in English and then only translated into other languages—a full and comprehensive Polish adaptation has not been created (Bergdahl et al., 2023). In parallel to here reported initiative, the short AI attitude scale AIAS-4 (AIAS-4-PL) adaptation has been performed by Juranek & Pochwatko (2026), the measurement was also translated to Polish in separate cross-cultural study (Talik et al., 2025).

Research conducted to date considering attitudes toward AI in the Polish context has been performed mostly in cross-cultural comparison settings. The measurements used included the earlier mentioned tools (short GAAIS, AIAS-4), custom designs fitting the needs of a particular project (such as the “Perceptions on AI by the Citizens of Europe” questionnaire, PAICE; Scantamburlo et al., 2024), and methods like qualitative interviews (Smola et al., 2025a). The studies showed that Poles, compared to other nations, are among the highest in trust and approval of AI (Bergdahl et al., 2023; Scantamburlo et al., 2024). This higher trust is affected, among other factors, by the positive influence of social norms (Smola et al., 2025b), with other significant predictors identified as gender, religious practices, religious development, relatedness, and competence (Bergdahl et al., 2023; Kozak & Fel, 2024).

To address the limited availability of psychometrically validated instruments for assessing general attitudes toward AI in the Polish language, we decided to translate, research factor structure, validate and make available to Polish research community and international researchers investigating Polish-language participants, the ATTARI-12 questionnaire (Stein et al., 2024).

The choice of the ATTARI-12 scale was motivated by the advantages it offers over previously developed instruments. First and foremost, none of the existing scales beside ATTARI-12 explicitly incorporate the features substantial to the concept of attitudes: the affective, behavioral, and cognitive facets (Stein et al., 2024). In addition, the ATAI exhibited low reliability (Sindermann et al., 2021; Stein et al., 2024). Furthermore, alternative measures often employ multi-factor structures; while this may offer more nuance and insight, it may simultaneously render the interpretation of results more complicated (GAAIS, ATAI; Schepman & Rodway, 2023; Sindermann et al., 2021), which was considered a disadvantage. The ATTARI-12 has multiple language versions that position it as a valuable instrument for cross-national and cross-cultural comparative studies. English and German versions were developed initially (Stein et al., 2024), followed by Korean (Lee et al., 2025) and Romanian adaptations (Maier et al., 2025). Additionally, a context-specific English version focusing on AI in work, healthcare, and education exists (Gnambs et al., 2025).

This paper aims to investigate the psychometric properties of the Polish version of ATTARI-12 (ATTARI-12-PL) and to uncover the relationships between general AI attitude, AI anxiety, negative attitudes toward robots and belief in human nature uniqueness while testing the adaptation’s congruent validity. Furthermore, the work attempts to determine the utility of the ATTARI-12-PL general factor by comparison of its similarity via correlation with general AI attitude measured by AIAS-4-PL, another popular and short, general AI attitude measuring tool created by Grassini (Grassini, 2023) and adapted earlier into Polish by the authors of this work.

The permission for the adaptation of the questionnaire had been obtained from the authors of the original ATTARI-12 measure (Stein et al., 2024). The hypotheses and planned data analysis steps were preregistered at https://aspredicted.org/LKZ_H97.

METHOD

Ethics statement

The research study received approval from the Ethics Committee (approval number: 11/IV/2024) of the Institute of Psychology of the Polish Academy of Sciences and was conducted in full accordance with the Declaration of Helsinki.

Participants

Sample of initial 494 participants had been recruited using research platforms (Prolific & SW Research “Ankieteo”) and social network adds. From the “Ankieteo” platform originated 328 participants, another 115 were recruited via social network adds, with remaining 51 participants coming from Prolific. In line with preregistered rules for excluding observations, records from participants who failed attention check question or for which any data point was missing, were excluded. The number of excluded participants due to failed attention check, was unexpectedly high, particularly on the “Ankieteo” platform, where it reached rate of 42%. In all, exclusion rate for all reasons (missing data and failed attention check) reached 32.8% of the overall responses gathered pool, leaving the number of participants which answers could be analyzed at N=332. The unexpectedly high number of drop-out data turned the available responses pool below the preregistered threshold of minimum 360 complete data records to proceed with calculation. However, the commonly recommended threshold allowing factor analysis was still exceeded by the number of collected responses, which stands at 300 observations and a minimum ratio of 10:1 responses per questionnaire item/variable (Comrey & Lee, 2013; Yong & Pearce, 2013). Considering carefully these criteria and recommendations, the decision was made to proceed with analysis of the collected data.

The Prolific participants had been compensated for their participation with the recommended Prolific hourly rate in British Pound, the respondents volunteering participation through remaining networks/platforms were

not compensated. The age distribution of participants is depicted in Table 1 and the demographic characteristics of the sample (N=332) are depicted in Table 2.

Table 1. Age Distribution of Participants.

	Mean	Median	SD	Min	Max
Age	37.95	34	15.58	18	78

Table 2. Demographic Characteristics of Participants.

	Count	%
Gender		
Male	118	35.5
Female	209	63.0
Other	4	1.2
Not Declared	1	0.3
Residence		
Rural	63	19.0
Small Town (= < 50 k)	62	18.7
Medium City (= < 500 k)	123	37.0
Large City (> 500 k)	83	25.3
Education		
Primary	7	2.1
Secondary	153	46.1
Tertiary: Bachelor	65	19.6
Tertiary: Masters	98	29.5
Tertiary: Doctorate	9	2.7

Measures

ATTARI-12-PL. The Polish adaptation of the ATTARI-12 questionnaire was developed through a multi-step process. Initially, a professional translator translated the original English version into Polish. Subsequently, independent translations were generated using three machine translation tools: Google Translate, DeepL, and Microsoft Copilot. Each item in every version was then back-translated to English using the reverse method. In cases where convergence between all the versions occurred, items were automatically picked to form the Polish questionnaire. For the cases where convergence did not occur, translations of items were chosen with the consultation and advice of Polish language specialist, specialized in teaching Polish as foreign language, having also knowledge of English on C1 level of the Common European Framework of Reference for Languages (CEFR), to ensure the translations convey the original meaning.

The original ATTARI-12 questionnaire is a 12 item scale developed in English and German versions by Stein and colleagues (2024) in set of three studies. Half of the items (6) have negative valence and are reverse-coded,

meaning that agreeing with the item statement represents negative attitude toward AI, example: “I have strong negative emotions about AI.”. Each facet of attitude (affective, cognitive, behavioral) is represented by four items, two with positive and two with negative valence. Participants respond using a 5-point answer format (1 – strongly disagree, 5 – strongly agree).

Artificial Intelligence Attitude Scale – Polish Version (AIAS-4-PL). The English version of the 4-item version of AIAS-4 (Grassini, 2023) has been translated, adapted and validated successfully in Polish context by the authors of this paper, with the psychometric properties of the scale to be published in another work. This brief scale measures attitude toward AI as one general factor. Responses are gathered using a 10-point scale (1 – strongly disagree to 10 – strongly agree).

The Negative Attitude Towards Robots Scale – Polish Version (NARS-PL) is an adaptation developed by Pochwatko et al., (2015) based on the original Japanese 14-item scale created by Nomura and colleagues (2006). The Polish version consists of 12 items, assessing negative attitudes toward robots across two subscales: NARHT (negative attitudes toward robots with human traits) and NATIR (negative attitude toward interaction with robots). Participants respond using a 7-point scale (1 – strongly disagree, 7 – strongly agree).

Belief in Human Nature Uniqueness Scale (BHNUS). Polish version (Pochwatko et al., 2015) of the scale that assesses the degree to which individuals uphold the idea of human nature as unique to their own group and reject the notion that robots could ever be fundamentally similar to humans. The original scale has been developed by Piçarra (2014). The scale consists of 6 items and is unidimensional. Participants respond on a 7-point scale (1 – totally disagree to 7 – totally agree). BHNU is considered as sociocognitive rather than fixed idiosyncratic personality factor (Giger et al., 2024). This belief has been shown to play an important role in shaping reactions to autonomous technologies, particularly in contexts where technology is perceived as approaching human-like capabilities or social roles. The BHNU score has been reported to correlate with the cognitive ability of understanding thoughts and feelings of others (perspective taking), religiosity, interest for science fiction, negative attitude toward robots, attribution of warmth traits in robots, emotional appraisal of robots and attitude toward development of robots with human traits (Giger et al., 2024). Strong BHNU has been shown to deepen the uncanny valley effect phenomenon (Masahiro et al., 2012; Ratajczyk et al., 2024). The BHNU provides a useful framework to test whether general attitudes toward AI are in any measure associated with identity-based, essentialism beliefs regarding the boundary between humans and intelligent technological systems, which association has been shown in the case of robotic technology.

Artificial Intelligence Anxiety Scale – Polish Version (AIAS PL). The Polish adaptation is based on a questionnaire developed by Wang & Wang (2022). The scale comprises 21 items that assess four facets of AI anxiety:

learning anxiety, AI configuration anxiety, job replacement anxiety and the sociotechnical blindness component. The Polish version has been validated by Fortuna et al., (2023; Modliński et al., 2023). The participants respond to 7-point scale (1 – totally disagree to 7 – totally agree).

Procedure

The study was conducted online using the Google Forms platform. Participants were informed that the study was anonymous and no personal data would be collected. They were assured they could withdraw at any stage without repercussions. After accepting an informed consent form and reading a brief instruction, participants responded to items on the measurement scales. One of the items included was attention check question: “To confirm that you are a real person and not a ‘bot,’ please select ‘strongly agree’ for this item.”. The average time for completion of the responses was below 7 minutes.

RESULTS

Descriptive statistics were acquired for the overall AI attitude score, each attitude facet (affective, cognitive and behavioral) and every single item of the questionnaire. The skewness values obtained (ranging from -0.37 to 0.21), are close to zero, indicating a relatively symmetric distribution. The kurtosis values (-1.25 to -0.62) display flatter, compared to normal, shape, exhibiting a platykurtic distribution. The descriptive statistics acquired are presented in Table 3.

The Cronbach alpha coefficient was employed to assess the reliability of the questionnaire. The entire scale exhibits a Cronbach’s alpha of 0.94, indicating strong coherence across the items, though such a high value on

a relatively short 12-item scale may also suggest some redundancy, with items potentially capturing overlapping aspects of the construct.

To assess data suitability for factor analysis, the Kaiser-Meyer-Olkin (KMO) measure and Bartlett’s test of sphericity were used. The KMO value of 0.93 indicates excellent sampling adequacy, and the significant Bartlett test [$\chi^2(66) = 2974.50, p < 0.001$] supports validity for factor analysis.

The investigation of the factor structure of the scale, as preregistered, has been conducted first with the confirmatory factor analysis (CFA). To assess the model fit, a number of fit indices were employed. Comparative Fit Index (CFI) result at 0.84 accompanied by the Tucker-Lewis Index (TLI) value of 0.80 suggested a room for refinement of the model. Root Mean Square Error of Approximation (RMSEA) of 0.16 was above the suggested cutoff of 0.06 (Xia & Yang, 2019), and the Standardized Root Mean Square Residual (SRMR) obtained was 0.07, additionally a significant result was obtained from the chi-square test for exact fit ($\chi^2(54) = 530.78, p < 0.001$).

The fit of the single factor model was poor. The initial CFA has shown that the dimensionality of the ATTARI-12-PL cannot be fully explained by a unidimensional structure. This outcome mirrored findings from the original scale development, suggesting the presence of systematic variance stemming from design features. Therefore an approach following steps of the authors of the original scale, utilizing the bifactorial analysis, has been employed. In their work Stein and colleagues (2024) at first tested a general single factor model, with no minor secondary orthogonal multidimensionality, obtaining results similar to those found in this work, concluding the

Table 3. Descriptive statistics for overall score, three attitude facets and every item of ATTARI-12-PL.

	Mean	Median	SD	Skewness	Kurtosis
Overall Score	3.10	3.08	0.96	-0.13	-0.62
Behavioral Facet	3.09	3.00	1.13	-0.06	-0.83
Cognitive Facet	3.12	3.25	0.91	-0.11	-0.21
Affective Facet	3.09	3.00	1.09	-0.19	-0.77
Item 1	2.79	3.00	1.19	0.11	-0.75
Item 2	3.34	3.00	1.31	-0.29	-1.05
Item 3	3.32	3.00	1.28	-0.37	-0.91
Item 4	3.08	3.00	1.23	-0.11	-0.84
Item 5	3.10	3.00	1.39	-0.12	-1.25
Item 6	3.37	3.00	1.13	-0.26	-0.63
Item 7	2.98	3.00	1.30	0.07	-1.08
Item 8	2.77	3.00	1.25	0.21	-0.99
Item 9	2.96	3.00	1.28	-0.05	-1.02
Item 10	3.23	3.00	1.13	-0.23	-0.62
Item 11	3.13	3.00	1.20	-0.27	-0.78
Item 12	3.12	3.00	1.30	-0.12	-1.11

assumption for a single factor not taking into account features such as facets of attitudes was too strong. Therefore, they tested three bifactorial models. First bifactorial model assumed a general factor plus two orthogonal specific factors considering attitude facets: cognitive or affective items. Following Eid and colleagues (2017) work regarding anomalous results in G-factor models, they excluded behavioral facet to consider it as reference. In the second bifactor model, negative wording (valence) of items had been selected as the specific orthogonal factor. Their third, final model included all three specific factors identified in the two previous bifactorial models and was characterized as the best fit. In this work all these three bifactor models had been tested, additionally a model not excluding behavioral facet as reference had also been tried. The fit indices and parameters obtained for each model, including the initial single factor model had been depicted in Table 4.

In case of this work the best fitting model, that also met the commonly established criteria of good fit, has been the bifactorial model with each of the three attitude facets plus the negative wording of the items considered as specific, orthogonal factors in addition to the common general factor. The CFI for this model is 0.994 and the TLI 0.988. The chi-square test for exact fit returned a result close to insignificance [$\chi^2(36) = 55.192$, $p = 0.021$]. The RMSEA obtained for this model is below the 0.06 threshold standing at 0.04, with SRMR at 0.021.

In addition to the factor analyses, and in accordance with the preregistration, a Principal Components Analysis (PCA) was conducted. PCA revealed that the first principal component accounted for 59% of the variance, with all items loading substantially (0.60–0.86). A parallel analysis supported a single dominant component. While these results suggest the presence of a strong general factor, they do not override the poor fit of the strict single-factor CFA model. Instead, they complement the bifactor CFA findings, which demonstrated that the ATTARI-12-PL is best represented by a strong general factor alongside systematic variance related to facets and item wording.

The validity of the questionnaire, as preregistered, was assessed by calculating Pearson correlations with scores from related measures (AIAS-4-PL, NARS-PL,

BHNUS, and AIAS-PL). Correlations had been calculated for overall scores and subscales for those questionnaires that have subscales. There have been high, positive and significant correlation of overall ATTARI-12-PL score with AIAS-4-PL score ($r = 0.85$, $p < 0.001$). Correlations with all the other measurements were significant and negative, as expected. The largest negative correlation was noticed between the AI attitude as measured by ATTARI-12-PL and the negative attitude toward interaction with robots (NATIR) subscale of NARS-PL ($r = -0.72$, $p < 0.001$). This is followed by second largest negative correlation with AI anxiety as measured by AIAS-PL ($r = -0.69$, $p < 0.001$) and its related subscales (AI configuration anxiety: $r = -0.64$; AI learning anxiety: $r = -0.63$; AI job loss anxiety: $r = -0.60$; all at $p < 0.001$ significance level), then negative attitude toward robots with human traits subscale of NARS-PL ($r = -0.56$, $p < 0.001$), AI social blindness factor of AIAS-PL ($r = -0.51$, $p < 0.001$) and least negative correlation with belief in human nature uniqueness ($r = -0.27$, $p < 0.001$). The correlation matrix revealed also a very strong positive correlation of AI anxiety with the NATIR subscale of NARS-PL ($r = 0.80$, $p < 0.001$). The complete correlation matrix is depicted in Table 5.

To explore potential relationships between demographic factors and attitudes toward AI, a multiple linear regression model was conducted with age, gender, place of residence size, and education level as predictors. Assumptions of homoscedasticity and normality were checked: the residuals versus fitted plot indicated no clear violations of homoscedasticity, and the Q-Q plot showed residuals roughly aligned with the reference line, suggesting approximate normality.

The model was statistically significant overall ($F = 2.46$, $p = 0.0058$) but explained only 7.8% of the variance in AI attitude scores ($R^2 = 0.078$), indicating limited explanatory power. Among the predictors, male gender (relative to female) showed a small but statistically significant positive association with AI attitude scores ($\beta = 0.340$, $p = 0.0021$). Doctoral education level approached significance with a positive coefficient ($p = 0.0635$). These results suggest that demographic factors account for only a minor portion of individual differences in attitudes toward AI.

Table 4. Goodness of fit for rivaling confirmatory factor models for the ATTARI-12-PL.

Model	Factors	χ^2	df	p	CFI	TFI	RMSEA	SRMR
Single								
	G	530.780	54	> 0.0001	0.839	0.803	0.163	0.073
Bifactor								
	G + wording	296.357	48	> 0.0001	0.916	0.885	0.125	0.045
	G + facets: a + c	413.584	46	> 0.0001	0.876	0.822	0.155	0.065
	G + wording + facets: a + c	150.414	40	> 0.0001	0.963	0.938	0.091	0.036
	G + wording + facets: a + b + c	55.192	36	0.021	0.994	0.988	0.040	0.021

Factors in model: G – general factor; wording – negative item wording factor; facets – tripartite attitude facets: a – affective, b – behavioral, c – cognitive.

Table 5. Correlation matrix of ATTARI-12-PL and its subscales with AIAS-4-PL, NARS-PL subscales, BHNUS and AIAS-PL and its subscales.

	ATT-12-PL Score	ATT-12-PL Cogn	ATT-12-PL Aff	ATT-12-PL Behav	AIAS-4-PL Score	NARS-PL NATIR	NARS-PL NARHT	BHNUS	AIAS PL Score	AIAS PL Learn	AIAS PL Job	AIAS PL Conf	AIAS PL Blind
ATT-12-PL Score													
ATT-12-PL Cogn	0.89 ***												
ATT-12-PL Aff	0.93 ***	0.75 ***											
ATT-12-PL Behav	0.93 ***	0.75 ***	0.80 ***										
AIAS-4-PL Score	0.85 ***	0.75 ***	0.79 ***	0.80 ***									
NARS-PL NATIR	-0.72 ***	-0.60 ***	-0.69 ***	-0.67 ***	-0.60 ***								
NARS-PL NARHT	-0.56 ***	-0.46 ***	-0.53 ***	-0.55 ***	-0.46 ***	0.71 ***							
BHNUS	-0.27 ***	-0.21 ***	-0.20 ***	-0.32 ***	-0.23 ***	0.47 ***	0.49 ***						
AIAS PL Score	-0.69 ***	-0.56 ***	-0.69 ***	-0.64 ***	-0.55 ***	0.80 ***	0.58 ***	0.38 ***					
AIAS PL Learn	-0.63 ***	-0.52 ***	-0.61 ***	-0.61 ***	-0.50 ***	0.73 ***	0.45 ***	0.37 ***	0.89 ***				
AIAS PL Job	-0.60 ***	-0.48 ***	-0.61 ***	-0.54 ***	-0.47 ***	0.68 ***	0.54 ***	0.29 ***	0.88 ***	0.61 ***			
AIAS PL Conf	-0.64 ***	-0.51 ***	-0.64 ***	-0.59 ***	-0.50 ***	0.78 ***	0.58 ***	0.41 ***	0.91 ***	0.77 ***	0.77 ***		
AIAS PL Blind	-0.51 ***	-0.41 ***	-0.54 ***	-0.46 ***	-0.40 ***	0.60 ***	0.53 ***	0.23 ***	0.79 ***	0.52 ***	0.78 ***	0.73 ***	

Note. ATT-12-PL Score: ATTARI-12-PL general score; ATT-12-PL Cogn: ATTARI-12-PL cognitive facet subscale; ATT-12-PL Aff: ATTARI-12-PL affective facet subscale; ATT-12-PL Behav: ATTARI-12-PL behavioral facet subscale; NARS-PL NATIR: negative attitude toward interaction with robots subscale; NARS-PL NARHT: negative attitude toward robots with human traits subscale; BHNUS: belief in human nature uniqueness scale; AIAS PL Learn: learning anxiety subscale; AIAS PL Job: job replacement anxiety subscale; AIAS PL Conf: AI configuration anxiety subscale; AIAS PL Blind: sociotechnical blindness subscale. *** $p < 0.001$.

Table 6. Multiple linear regression model summarizing the relationships between age, gender, the size of the place of residence and education as predictors of ATTARI-12-PL total score.

	Estimate	SE	p	Significance
Intercept	3.1394	0.3869	$p < 0.001$	***
Age	-0.0058	0.0035	0.1022	
Gender (Reference: Female)				
Male	0.3403	0.1095	0.0021	**
Other	-0.5589	0.4809	0.2460	
Prefer not to say	-1.4186	0.9477	0.1354	
Residence (Reference: Rural)				
Small Town	0.1075	0.1720	0.5324	
Medium City	-0.1060	0.1482	0.4750	
Large City	-0.1003	0.1626	0.5378	
Education (Reference: Primary)				
Secondary	0.0802	0.3622	0.8250	
Tertiary: BA	0.2370	0.3747	0.5275	
Tertiary: MA	0.0216	0.3689	0.9533	
Tertiary: DR	0.8926	0.4795	0.0636	.

Note. SE: standard error; BA: Bachelor's degree; MA: Master's degree; DR: Doctoral degree. *** $p < 0.001$, ** $p < 0.01$.

DISCUSSION

The study not only adopted and investigated the psychometric properties of the Polish version of the ATTARI-12 questionnaire but also explored significant associations between the tripartite (affective, behavioral, and cognitive) AI attitudes and attitudes toward humanoid robots, belief in human nature's uniqueness, and fear of AI.

The poor fit of the confirmatory single-factor model exhibited more complex factorial structure of the adopted measurement, beyond pure unidimensional construct. While the ATTARI-12-PL strongly loaded on a general factor, it also contained additional systematic variance. This variance reflects the tripartite facets of attitudes (affective, cognitive, behavioral) and the influence of item valence (positive vs. negative wording). Our results show that nuanced bifactor modeling is the most appropriate analytic framework for this instrument.

This research showed that the AI attitude had the strongest inverse relationship with the negative attitude toward interacting with robots (NATIR). This negative relationship was even stronger than the negative association between AI attitude scores and AI anxiety (both general and specific types), as well as the negative attitude toward robots with human traits. Relationships in the same direction but to a lesser extent were observed between the aforementioned dimensions and the AIAS-4-PL measure.

This result suggests that AI attitude measures can be considered moderate proxies for standpoints toward robots. What is interesting though, is that AI attitude score is only mildly negatively associated with belief in human nature uniqueness. This result can be viewed as showing belief in human nature uniqueness possibly does not involve very strong opposition and negative stance toward general AI, which might come from the fact that AI systems come in much more varieties and AI is a more vague category than the concept of humanoid robots, which additionally directly compete with humans in their appearance. Of course, AI systems and components are technically crucial to building the more advanced humanoid robots; however, AI is already present and embedded in a multitude of commonly, everyday used applications and machines, and the awareness of this among respondents is possibly reflected in the results presented here.

Our results also showed a strong, positive relationship between AI anxiety, as measured by the AIAS-PL, and the NATIR component of the NARS-PL scale. This finding suggests that a negative attitude toward interacting with robots might be an indicator of AI anxiety, and conversely, AI anxiety could be a strong predictor of reluctance to interact with robots.

This investigation exhibited also a strong positive relationship of the researched and adapted scale, which itself is quite brief, with very short AI attitude measurement scale which is AIAS-4-PL. This result might be

useful for consideration of tools by researchers when brevity and little number of items is crucial. Although ATTARI-12-PL is more elaborate scale, including negative-formulated items and respecting the tripartite character of attitudes, the significant, positive correlation of 0.85 exhibits also high utility of the shorter questionnaire.

Limitations

The current research has limitations. The model fit was acceptable, but only after trying multiple analyses. The research utilized convenience sampling and the sample is not representative. Many sample parameters deviated from national proportions, for example women were overrepresented (63% female vs. 35.5% male). Moreover, as participants were recruited through online platforms (Prolific, Ankieto, and social networks), bias toward more digitally literate populations may have been introduced, and relatively high exclusion rate may have reduced diversity. This non-representativeness poses challenges to the generalizability of the findings, although the minimal skewness and kurtosis in the data regarding the AI attitude score, suggests that the obtained distribution of scores is fairly symmetrical and close to normal distribution.

A further major limitation concerns data quality issues. We experienced approximately 42% of attention check response failure from one of the recruitment sources, and relatively low average completion time, under 7 minutes for a survey exceeding 50 items. This points to possibility of encountering careless responders and “speeders”, who might have introduced substantial noise into the dataset. Although mitigation steps were taken, the possible contamination reduces confidence of the findings and the exclusion itself might have biased the results as retention of only attentive and motivated responders could result in “self-selection” problem.

To mitigate the sample and data quality issues, future research should aim to gather larger, more representative samples, with utilization of methods like stratified or probability-based sampling. Recruitment should be conducted across a wider range of demographic categories (including older adults, rural residents, and individuals with lower digital engagement), potentially employing survey methods beyond online channels. Additionally, more layers of quality assurance should be introduced, including more sophisticated and multiple attention checks (e.g., instructional manipulation checks, consistency checks across related items, and bogus items), minimum completion time thresholds deduced from pilot-tested durations, and the use of bot-detection mechanisms on survey platforms.

Another limitation of the study pertains to the test of validity of the questionnaire, our study did not include any measures to assess discriminant validity, which would demonstrate the lack of correlation between it and unrelated measures.

Finally, the results obtained regarding the relationship of AI attitude and AI anxiety, attitudes toward robots and belief in human nature uniqueness, although insightful, are

of only correlational character. It would be beneficial to use the results as starting point to directed investigation that could reveal the casual relationship and potential latent factors that underly and determine both the attitudes toward robots and AI in whole variety of their facets.

CONCLUSION

The ATTARI-12-PL is a novel, brief and compact tool aiming to quickly and reliably measure AI attitude among people. This kind of comprehensive tool, encompassing negative-worded items and built with respect of the concept of tripartite attitude components in design was not available until now for researchers and psychologists operating within the Polish language research context. Our work exhibits that the Polish adaptation of the scale is reliable and valid and can be a valuable asset for prospect research in a contemporary world characterized by the rapid advancement of artificial intelligence.

FUNDING

This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

DISCLOSURE STATEMENT

The authors report there are no competing interests to declare.

INFORMED CONSENT STATEMENT

All participants provided informed consent prior to their inclusion in the study. The consent process ensured that participants were fully aware of the study's purpose, procedures, potential risks, and benefits. Participants were informed that their participation was voluntary and that they could withdraw from the study at any time without any consequences. Confidentiality and anonymity of the participants were maintained throughout the study.

PATIENT INFORMED CONSENT STATEMENT: N/A

ETHICAL APPROVAL

This research study received approval from the Research Ethics Committee at the Institute of Psychology, Polish Academy of Sciences (approval no. 11/IV/2024) and was conducted in full accordance with the Declaration of Helsinki.

DATA AVAILABILITY STATEMENT

The data supporting the results or analyses presented in this paper are available in the OSF repository at the following link: https://osf.io/vneyh/overview?view_only=5542c10ac3ab4e90b90a7a792542ff7b

REFERENCES

- Ajzen, I. (2001). Nature and Operation of Attitudes. *Annual Review of Psychology*, 52(1), 27–58. <https://doi.org/10.1146/annurev.psych.52.1.27>
- Bergdahl, J., Latikka, R., Celuch, M., Savolainen, I., Mantere, E. S., Savela, N., & Oksanen, A. (2023). Self-determination and attitudes toward artificial intelligence: Cross-national and longitudinal perspectives. *Telematics and Informatics*, 82, 102013. <https://doi.org/10.1016/j.tele.2023.102013>
- Bharadiya, J. (2023). Artificial intelligence in transportation systems a critical review. *American Journal of Computing and Engineering*, 6(1), 34–45. <https://doi.org/10.47672/ajce.1487>
- Blauth, T. F., Gstrein, O. J., & Zwitter, A. (2022). Artificial intelligence crime: An overview of malicious use and abuse of AI. *Ieee Access*, 10, 77110–77122. <https://doi.org/10.1109/access.2022.3191790>
- Bohner, G., & Dickel, N. (2011). Attitudes and attitude change. *Annual Review of Psychology*, 62, 391–417. <https://doi.org/10.1146/annurev.psych.121208.131609>
- Brewer, P. R., Bingaman, J., Paintsil, A., Wilson, D. C., & Dawson, W. (2022). Media Use, Interpersonal Communication, and Attitudes Toward Artificial Intelligence. *Science Communication*, 44(5), 559–592. <https://doi.org/10.1177/10755470221130307>
- Briñol, P., Petty, R. E., & Stavrak, M. (2019). Structure and function of attitudes. *Social Psychology*. <https://doi.org/10.1093/acrefore/9780190236557.013.320>
- Cave, S., & Dihal, K. (2019). Hopes and fears for intelligent machines in fiction and reality. *Nature Machine Intelligence*, 1(2), 74–78. <https://doi.org/10.1038/s42256-019-0020-9>
- Chen, L., Chen, P., & Lin, Z. (2020). Artificial intelligence in education: A review. *Ieee Access*, 8, 75264–75278. <https://doi.org/10.1109/access.2020.2988510>
- Chudziak, J., Gambin, T., Gawrysiak, P., & Muraszewicz, M. (2022). O sztucznej inteligencji. In *Sztuczna inteligencja dla inżynierów. Metody ogólne*. (Maruszkiewicz Mieczysław, Nowak Robert, pp. 1–19). Oficyna Wydawnicza Politechniki Warszawskiej.
- Comrey, A. L., & Lee, H. B. (2013). *A first course in factor analysis*. Psychology press. <https://doi.org/10.4324/9781315827506>
- Eid, M., Geiser, C., Koch, T., & Heene, M. (2017). Anomalous results in G-factor models: Explanations and alternatives. *Psychological Methods*, 22(3), 541. <https://doi.org/10.1037/met0000083>
- Fishbein, M., & Ajzen, I. (1977). *Belief, attitude, intention, and behavior: An introduction to theory and research*. <https://doi.org/10.2307/2065853>
- Fortuna, P., Łysiak, M., Chumak, M., & McNeill, M. (2023). Barriers of human and nonhuman agents' integration in positive hybrid systems: the relationship between the anthropocentrism, artificial intelligence anxiety, and attitudes towards humanoid robots. *Journal for Perspectives of Economic Political and Social Integration*, 28(2), 121–149.
- Giger, J.-C., Piçarra, N., Pochwatko, G., Almeida, N., Almeida, A. S., & Costa, N. (2024). Development of the Beliefs in Human Nature Uniqueness Scale and Its Associations With Perception of Social Robots. *Human Behavior and Emerging Technologies*, 2024(1), 5569587. <https://doi.org/10.1155/2024/5569587>
- Glasman, L. R., & Albarracín, D. (2006). Forming attitudes that predict future behavior: A meta-analysis of the attitude-behavior relation. *Psychological Bulletin*, 132(5), 778. <https://doi.org/10.1037/0033-2909.132.5.778>
- Gnambs, T., Stein, J.-P., Appel, M., Griese, F., & Zinn, S. (2025). An economical measure of attitudes towards artificial intelligence in work, healthcare, and education (ATTARI-WHE). *Computers in Human Behavior: Artificial Humans*, 3, 100106. <https://doi.org/10.1016/j.chbah.2024.100106>
- Grassini, S. (2023). Development and validation of the AI attitude scale (AIAS-4): A brief measure of general attitude toward artificial intelligence. *Frontiers in Psychology*, 14, 1191628. <https://doi.org/10.3389/fpsyg.2023.1191628>
- Guingrich, R., & Graziano, M. (2024). P (doom) Versus AI Optimism: Attitudes Toward Artificial Intelligence and the Factors that Shape Them. *PsyArXiv*. <https://doi.org/10.31234/osf.io/cenwv>
- Holzinger, A., Fister Jr, I., Fister, I., Kaul, H.-P., & Asseng, S. (2024). Human-Centered AI in smart farming: Towards Agriculture 5.0. *IEEE Access*. <https://doi.org/10.1109/access.2024.3395532>
- Jáuregui-Velarde, R., Andrade-Arenas, L., Molina-Velarde, P., & Yac-tayo-Arias, C. (2024). Financial revolution: A systemic analysis of artificial intelligence and machine learning in the banking sector. *International Journal of Electrical & Computer Engineering* (2088-8708), 14(1). <https://doi.org/10.11591/ijece.v14i1.pp1079-1090>
- Juranek, J. F. (2025). Multimodal Italian Sign Language Recognition with Radar-Video Late Fusion on the MultiMeDaLIS Dataset. 2025 IEEE/CVF International Conference on Computer Vision Workshops (ICCVW) 5079–5085. <https://doi.org/10.1109/ICCVW69036.2025.00528>
- Juranek, J. F. (2026). Multimodal Avatars and Human Simulations in Teaching Practical Skills and Training of Clinical Psychology, Psychotherapy, and Related Mental Health Specialists. *Psychology Learning & Teaching*, 25(1), 5–19. <https://doi.org/10.1177/14757257261426784>
- Juranek, J. F., & Pochwatko, G. (2026). The Polish adaptation of the Artificial Intelligence Attitude Scale: AIAS-4-PL. *Current Psychology*, 45(4), 410. <https://doi.org/10.1007/s12144-025-08930-5>
- Kaya, F., Aydin, F., Schepman, A., Rodway, P., Yetişensoy, O., & Demir Kaya, M. (2022). The Roles of Personality Traits, AI Anxiety, and Demographic Factors in Attitudes toward Artificial Intelligence. *International Journal of Human–Computer Interaction*, 1–18. <https://doi.org/10.1080/10447318.2022.2151730>
- Khogali, H. O., & Mekid, S. (2023). The blended future of automation and AI: Examining some long-term societal and ethical impact features. *Technology in Society*, 73, 102232. <https://doi.org/10.1016/j.techsoc.2023.102232>
- Kieslich, K., Lünich, M., & Marcinkowski, F. (2021). The Threats of Artificial Intelligence Scale (TAI) Development, Measurement and Test Over Three Application Domains. *International Journal of Social Robotics*, 13, 1563–1577. <https://doi.org/10.1007/s12369-020-00734-w>
- Kordzadeh, N., & Ghasemaghaei, M. (2022). Algorithmic bias: Review, synthesis, and future research directions. *European Journal of Information Systems*, 31(3), 388–409. <https://doi.org/10.1080/0960085X.2021.1927212>
- Kozak, J., & Fel, S. (2024). How sociodemographic factors relate to trust in artificial intelligence among students in Poland and the United Kingdom. *Scientific Reports*, 14(1), 28776. <https://doi.org/10.1038/s41598-024-80305-5>
- Kubassova, O., Shaikh, F., Melus, C., & Mahler, M. (2021). History, current status, and future directions of artificial intelligence. *Precision Medicine and Artificial Intelligence*, 1–38. <https://doi.org/10.1016/b978-0-12-820239-5.00002-4>
- Lee, K., Cooper, A. F., & Grimmelmann, J. (2024). Talkin' 'Bout AI Generation: Copyright and the Generative-AI Supply Chain (The Short Version). *Proceedings of the Symposium on Computer Science and Law*, 48–63. <https://doi.org/10.1145/3614407.3643696>
- Lee, S. Y., Kim, H. M., & Ahn, H. (2025). Validation of the Korean Version of the Attitudes Toward Artificial Intelligence Scale (ATTARI-12). *Korean Journal of Culture and Social Issue*, 31(2), 269–292.
- Luttrell, A., & Sawicki, V. (2020). Attitude strength: Distinguishing predictors versus defining features. *Social and Personality Psychology Compass*, 14(8), e12555. <https://doi.org/10.1111/spc3.12555>
- Maier, R., Ghiocanu, Ștefania C., Costin, A., & Simion, I. (2025). Students' Attitude Towards Artificial Intelligence and Its Predictive Factors. *BRAIN. Broad Research in Artificial Intelligence and Neuroscience*, 16(1 Sup1), 264–275. <https://doi.org/10.70594/brain/16.s1/22>
- Masahiro, M., MacDorman, K. F., & Kageki, N. (2012). The uncanny valley. *IEEE Robotics & Automation Magazine*, 19(2), 98–100.
- Méndez-Suárez, M., Delbello, L., De Vega De Unceta, A., & Ortega Larrea, A. L. (2024). Factors Affecting Consumers' Attitudes Towards Artificial Intelligence. *Journal of Promotion Management*, 30(7), 1141–1158. <https://doi.org/10.1080/10496491.2024.2367203>

- Modliński, A., Fortuna, P., & Rożnowski, B. (2023). Robots onboard? Investigating what individual predispositions and attitudes influence the reactions of museums' employees towards the adoption of social robots. *Museum Management and Curatorship*, 1–25. <https://doi.org/10.1080/09647775.2023.2235678>
- Nomura, T., Suzuki, T., Kanda, T., & Kato, K. (2006). Measurement of negative attitudes toward robots. *Interaction Studies. Social Behaviour and Communication in Biological and Artificial Systems*, 7(3), 437–454. <https://doi.org/10.1075/is.7.3.14nom>
- Ostrom, T. M. (1969). The relationship between the affective, behavioral, and cognitive components of attitude. *Journal of Experimental Social Psychology*, 5(1), 12–30. [https://doi.org/10.1016/0022-1031\(69\)90003-1](https://doi.org/10.1016/0022-1031(69)90003-1)
- Park, J., & Woo, S. E. (2022). Who Likes Artificial Intelligence? Personality Predictors of Attitudes toward Artificial Intelligence. *The Journal of Psychology*, 156(1), 68–94. <https://doi.org/10.1080/00223980.2021.2012109>
- Piçarra, N. J. G. (2014). *Predicting intention to work with social robots*. <https://sapientia.ualg.pt/handle/10400.1/6854>
- Pochwatko, G., Giger, J.-C., Różańska-Walczyk, M., Świdrak, J., Kukiłka, K., Możaryn, J., & Piçarra, N. (2015). Polish version of the negative attitude toward robots scale (NARS-PL). *Journal of Automation Mobile Robotics and Intelligent Systems*, 9. https://doi.org/10.14313/jamris_2-2015/25
- Ratajczyk, D., Dakowski, J., & Łupkowski, P. (2024). The Importance of Beliefs in Human Nature Uniqueness for Uncanny Valley in Virtual Reality and On-Screen. *International Journal of Human–Computer Interaction*, 40(12), 3081–3091. <https://doi.org/10.1080/10447318.2023.2179216>
- Rice, R. E., & Wu, M.-Y. (2025). Difference in and Influences on Public Opinion About Artificial Intelligence in 20 Economies: Reducing Uncertainty Through Awareness, Knowledge, and Trust. *International Journal of Communication*, 19, 26.
- Scantamburlo, T., Cortés, A., Foffano, F., Barrué, C., Distefano, V., Pham, L., & Fabris, A. (2024). Artificial intelligence across europe: A study on awareness, attitude and trust. *IEEE Transactions on Artificial Intelligence*. <https://doi.org/10.1109/tai.2024.3461633>
- Schaur, M., & Matausch-Mahr, K. (2025). Assistive technology using artificial intelligence in the long-term care sector for persons with disabilities: A systematic literature review. *Technology and Disability*, 37(3), 259–269. <https://doi.org/10.1177/10554181251355420>
- Schepman, A., & Rodway, P. (2020). Initial validation of the general attitudes towards Artificial Intelligence Scale. *Computers in Human Behavior Reports*, 1, 100014.
- Schepman, A., & Rodway, P. (2023). The General Attitudes towards Artificial Intelligence Scale (GAAIS): Confirmatory Validation and Associations with Personality, Corporate Distrust, and General Trust. *International Journal of Human–Computer Interaction*, 39(13), 2724–2741. <https://doi.org/10.1016/j.chbr.2020.100014>
- Secinaro, S., Calandra, D., Secinaro, A., Muthurangu, V., & Biancone, P. (2021). The role of artificial intelligence in healthcare: A structured literature review. *BMC Medical Informatics and Decision Making*, 21(1), 125. <https://doi.org/10.1186/s12911-021-01488-9>
- Sindermann, C., Sha, P., Zhou, M., Wernicke, J., Schmitt, H. S., Li, M., Sariyska, R., Stavrou, M., Becker, B., & Montag, C. (2021). Assessing the Attitude Towards Artificial Intelligence: Introduction of a Short Measure in German, Chinese, and English Language. *KI - Künstliche Intelligenz*, 35(1), 109–118. <https://doi.org/10.1007/s13218-020-00689-0>
- Smola, P., Młodziński, I., Wojcieszko, M., Zwierczyk, U., Kobryn, M., Rzepecka, E., & Duplaga, M. (2025). Attitudes toward artificial intelligence and robots in healthcare in the general population: A qualitative study. *Frontiers in Digital Health*, 7, 1458685. <https://doi.org/10.3389/fdgth.2025.1458685>
- Stănescu, D. F., & Romaşcanu, M. C. (2024). The influence of AI Anxiety and Neuroticism in Attitudes toward Artificial Intelligence. *European Journal of Sustainable Development*, 13(4), 191–191. <https://doi.org/10.14207/ejsd.2024.v13n4p191>
- Stein, J.-P., Liebold, B., & Ohler, P. (2019). Stay back, clever thing! Linking situational control and human uniqueness concerns to the aversion against autonomous technology. *Computers in Human Behavior*, 95, 73–82. <https://doi.org/10.1016/j.chb.2019.01.021>
- Stein, J.-P., Messingschlager, T., Gnambs, T., Huttmacher, F., & Appel, M. (2024). Attitudes towards AI: Measurement and associations with personality. *Scientific Reports*, 14(1), 2909. <https://doi.org/10.1038/s41598-024-53335-2>
- Svenningsson, J., Höst, G., Hultén, M., & Hallström, J. (2022). Students' attitudes toward technology: Exploring the relationship among affective, cognitive and behavioral components of the attitude construct. *International Journal of Technology and Design Education*, 32(3), 1531–1551. <https://doi.org/10.1007/s10798-021-09657-7>
- Talik, W., Talik, E. B., & Grassini, S. (2025). Measurement invariance of the artificial intelligence attitude scale (AIAS-4): Cross-cultural studies in Poland, the USA, and the UK. *Current Psychology*, 44(19), 15758–15766. <https://doi.org/10.1007/s12144-025-08348-z>
- Wang, Y.-Y., & Wang, Y.-S. (2022). Development and validation of an artificial intelligence anxiety scale: An initial application in predicting motivated learning behavior. *Interactive Learning Environments*, 30(4), 619–634. <https://doi.org/10.1080/10494820.2019.1674887>
- Watts, T. F., & Bode, I. (2024). Machine guardians: The Terminator, AI narratives and US regulatory discourse on lethal autonomous weapons systems. *Cooperation and Conflict*, 59(1), 107–128. <https://doi.org/10.1177/00108367231198155>
- Xia, Y., & Yang, Y. (2019). RMSEA, CFI, and TLI in structural equation modeling with ordered categorical data: The story they tell depends on the estimation methods. *Behavior Research Methods*, 51(1), 409–428. <https://doi.org/10.3758/s13428-018-1055-2>
- Xu, Y., Liu, X., Cao, X., Huang, C., Liu, E., Qian, S., Liu, X., Wu, Y., Dong, F., & Qiu, C.-W. (2021). Artificial intelligence: A powerful paradigm for scientific research. *The Innovation*, 2(4). <https://doi.org/10.1016/j.xinn.2021.100179>
- Yong, A. G., & Pearce, S. (2013). A beginner's guide to factor analysis: Focusing on exploratory factor analysis. *Tutorials in Quantitative Methods for Psychology*, 9(2), 79–94. <https://doi.org/10.20982/tqmp.09.2.p079>

APPENDIX: A

ATTARI-12-PL

Instrukcja

W poniższym kwestionariuszu jesteście zainteresowani Twoim stosunkiem wobec sztucznej inteligencji. Sztuczna inteligencja może wykonywać zadania, które zazwyczaj wymagają ludzkiej inteligencji. Umożliwia maszynom wyczuwanie, działanie, uczenie się i adaptację w autonomiczny, podobny do ludzkiego sposób. Sztuczna inteligencja może być częścią komputera lub platformy internetowej, ale można ją również spotkać w różnych innych urządzeniach, takich jak roboty.

	1 Zdecydowanie się nie zgadzam	2 Nie zgadzam się	3 Ani się zgadzam, ani nie zgadzam	4 Zgadzam się	5 Zdecydowanie się zgadzam
1. Sztuczna inteligencja uczyni ten świat lepszym miejscem.					
2. Mam silne negatywne emocje w stosunku do sztucznej inteligencji.					
3. Chcę korzystać z technologii opartych na sztucznej inteligencji.					
4. Sztuczna inteligencja ma więcej wad niż zalet.					
5. Z niecierpliwością czekam na przyszły rozwój sztucznej inteligencji.					
6. Sztuczna inteligencja oferuje rozwiązania wielu światowych problemów.					
7. Preferuję technologie, które nie zawierają sztucznej inteligencji.					
8. Obawiam się sztucznej inteligencji.					
9. Wybrałbym/wybrałabym raczej technologię ze sztuczną inteligencją niż bez niej.					
10. Sztuczna inteligencja raczej tworzy problemy niż je rozwiązuje.					
11. Kiedy myślę o sztucznej inteligencji, mam głównie pozytywne odczucia.					
12. Raczej unikałbym/unikałabym technologii opartych na sztucznej inteligencji.					

Komponenty: poznawczy (pozycje 1, 4, 6, 10), afektywny (pozycje 2, 5, 8, 11), behawioralny (pozycje 3, 7, 9, 12). Pozycje o negatywnej walencji i odwróconym kodowaniu: 2, 4, 7, 8, 10, 12.