

ORIGINAL PAPER ΕΡΕΥΝΗΤΙΚΗ ΕΡΓΑΣΙΑ

Short Trust in Automation Scale Translation and validation in Greek

OBJECTIVE To translate and validate the Short Trust in Automation Scale (S-TIAS) in a sample of the general population in Greece. **METHOD** We applied the forward-backward translation method to adapt the S-TIAS into Greek. Reliability was assessed by calculating Cronbach's alpha for internal consistency and conducting a test-retest procedure to determine stability using the intraclass correlation coefficient. Construct validity was evaluated through confirmatory factor analysis, while concurrent validity was examined by correlating S-TIAS scores with those from the Artificial Intelligence Attitude Scale (AIAS-4). **RESULTS** We found that the S-TIAS had very good reliability since the intraclass correlation coefficient for the test-retest was 0.962 (95% confidence interval [CI]=0.915 to 0.983, $p<0.001$). Moreover, Cronbach's coefficient alpha for the S-TIAS was 0.912. We found that the Greek version of the S-TIAS had a one-factor structure as the original version. All indices indicated an acceptable one-factor model. In particular, RMSEA was <0.001 , GFI was 1.000, IFI was 1.000, NFI was 1.000, and CFI was 1.000. Concurrent validity of the Greek version of the S-TIAS was very good since we found statistically significant correlation between the S-TIAS and the AIAS-4 ($r=0.608$, $p<0.001$). **CONCLUSIONS** The Greek version of the S-TIAS is a reliable and valid tool to measure trust towards artificial intelligence in the general population.

ARCHIVES OF HELLENIC MEDICINE 2026, 43(6):787–791
ΑΡΧΕΙΑ ΕΛΛΗΝΙΚΗΣ ΙΑΤΡΙΚΗΣ 2026, 43(6):787–791

A. Katsiroumpa,¹
I. Moisoglou,²
P. Gallos,³
O. Galani,³
M. Tsiachri,¹
P. Lialiou,⁴
O. Konstantakopoulou,¹
K. Lamprakopoulou,⁵
P. Galanis¹

¹Laboratory of Clinical Epidemiology,
Faculty of Nursing, National and
Kapodistrian University of Athens,
Athens

²Faculty of Nursing, University of
Thessaly, Larissa

³Faculty of Nursing, University of West
Attica, Athens

⁴Research Laboratory of Computational
Biomedicine, Department of Digital
Systems, University of Pireus, Pireus

⁵Laboratory of Embedded Systems
Design and Applications, Department
of Electrical and Computer Engineering,
University of Peloponnese, Patras, Greece

Κλίμακα “Short Trust in Automation
Scale”: Μετάφραση και στάθμιση
στην ελληνική γλώσσα

Περίληψη στο τέλος του άρθρου

Key words

Artificial intelligence
Greek
Reliability
Short Trust in Automation Scale
Validity

Submitted 11.7.2025

Accepted 10.8.2025

Artificial intelligence (AI) is reshaping economies and public services worldwide, from clinical decision support and diagnostics to education, finance, and transportation. Yet, the societal benefits of AI hinge not only on technical performance but also on the extent to which people trust AI systems; trust that they will be safe, fair, transparent,

and accountable. International frameworks emphasize precisely these qualities: the Organisation for Economic Co-operation and Development (OECD) AI principles –the first intergovernmental standard on AI– call for inclusive growth, human-rights alignment, transparency, robustness, and accountability, setting a global baseline for

“trustworthy AI”. In parallel, UNESCO’s recommendation on the ethics of AI and the World Health Organization’s (WHO) guidance for AI in health operationalize ethical safeguards across domains, reinforcing requirements such as human oversight, privacy, and equity in high-stakes settings like healthcare.^{1–4}

Europe’s regulatory landscape further elevates trust as a prerequisite for adoption. The EU Artificial Intelligence Act introduces a risk-based legal framework with phased application, including early prohibitions on specific practices and graduated obligations for high-risk systems and general-purpose AI. As these obligations take effect across member states, including Greece, the demand grows for reliable, validated instruments that can measure public and professional trust to inform governance, education, and implementation strategies.

Empirically, trust in automation and AI has a long research lineage. Evidence consistently shows that trust is multi-determinant (human, system, and environmental factors) and predicts reliance behaviors; they also highlight persistent measurement inconsistency across studies. In organizational science, a comprehensive review of human trust in AI further documents how design features (e.g., transparency, reliability) and representational cues (e.g., embodiment, anthropomorphism) shape cognitive and affective facets of trust, while calling for validated, context-appropriate scales. Against this backdrop, public sentiment remains heterogeneous across geographies and populations, recent European and global surveys reveal both optimism about AI’s societal value and enduring concerns about risks, accountability, and labor market effects.^{5–9}

For measurement, the field has widely relied on the Trust between People and Automation Scale¹⁰ and related instruments; however, questions about dimensionality and validity in modern AI contexts persist. Recent psychometric work addressed these gaps by validating the 12-item Trust in Automation Scale (TIAS) across diverse AI applications and developing a three-item Short Trust in Automation Scale (S-TIAS)¹¹ with strong sensitivity and predictive validity. Additional contemporary evaluations similarly stress the need to validate trust questionnaires specifically for AI and to distinguish trust from distrust as related but separable constructs. Collectively, this evidence base supports the value of brief, psychometrically sound measures that can be deployed repeatedly in real-world, policy-relevant settings.

Despite progress internationally, validated Greek-language instruments for measuring trust in AI remain scarce. This gap limits surveillance of public and professional attitudes, constrains the evaluation of education or imple-

mentation interventions, and complicates cross-cultural comparison with studies using the original TIAS/S-TIAS or other validated tools.¹¹ In a policy environment shaped by the broader ethics and standards ecosystem, a Greek-validated short trust scale would enable locally interpretable and globally comparable evidence to guide adoption strategies in sectors such as healthcare, public administration, and education.

In this context, we translated and validated the Short Trust in Automation Scale¹¹ in a sample of the general population in Greece.

MATERIAL AND METHOD

The study included 95 participants from Greece, with data collection conducted in September 2025. The S-TIAS was translated and adapted into Greek using the forward-backward translation method.¹² Specifically, two independent scholars translated the original English version into Greek, followed by back-translation into English by two additional scholars. A fifth scholar reviewed the entire process and resolved any discrepancies to ensure semantic and conceptual equivalence.

The S-TIAS measures trust in AI, and includes three items with answers on a scale from 1 to 7. Total score is the mean score of the three items, and, thus, total S-TIAS score ranges from 1 to 7. Higher scores indicate higher levels of trust in AI.¹¹

We examined the reliability of the S-TIAS by calculating Cronbach’s alpha. Cronbach’s alpha higher than 0.6 indicates acceptable internal reliability. Also, we performed a test-retest study to examine the reliability of the S-TIAS by calculating the intraclass correlation coefficient (absolute agreement method).

We examined the construct validity of the S-TIAS by performing confirmatory factor analysis.¹³ We examined the concurrent validity of the S-TIAS using the Artificial Intelligence Attitude Scale (AIAS-4).¹⁴

The AIAS-4 measures general attitude toward AI and includes four items with answers on a scale from 1 to 10. Total score is the mean score of the 10 items, and, thus, total AIAS-4 score ranges from 1 to 10. Higher scores indicate more positive attitudes towards AI.¹⁴

Ethical considerations

We applied the guidelines of the Declaration of Helsinki to perform this study.¹⁵ Additionally, the study protocol was approved by the Ethics Committee of Faculty of Nursing, National and Kapodistrian University of Athens (protocol approval #01, September 14, 2025).

Statistical analysis

We performed confirmatory factor analysis (CFA) to examine

the construct validity of the S-TIAS. In particular, we calculated Chi-square/degree of freedom (χ^2/df); root mean square error of approximation (RMSEA); goodness of fit index (GFI); adjusted goodness of fit index (AGFI); Tucker-Lewis index (TLI); incremental fit index (IFI); normed fit index (NFI); comparative fit index (CFI).^{16,17} Acceptable value for χ^2/df is <5 , for RMSEA is <0.10 , and for all other measures in the CFA >0.90 . We used the AMOS version 21 (Amos Development Corporation, 2018) to conduct the CFA.

We calculated Pearson's correlation coefficient between the S-TIAS and the AIAS-4 to examine the concurrent validity of the S-TIAS. Also, we calculated intraclass correlation coefficients between the two S-TIAS measurements in test-retest study.

P-values less than 0.05 were considered as statistically significant. We used the Statistical Package for Social Sciences (SPSS) (IBM Corp released 2021, IBM SPSS Statistics for Windows, IBM Corp, Armonk, NY), version 28.0 for the analysis.

RESULTS

Study population included 105 participants from the general population. Among them, 83.3% ($n=75$) were females and 16.7% ($n=15$) were males. The mean age of our sample was 35.0 years (standard deviation: 14.5).

We found that the S-TIAS had very good reliability since the intraclass correlation coefficient for the test-retest was 0.962 (95% confidence interval [CI]=0.915 to 0.983, $p<0.001$). Moreover, Cronbach's coefficient alpha for the S-TIAS was 0.912.

Furthermore, the Greek version of the S-TIAS was observed to have a one-factor structure as the original version (fig. 1). All indices indicated an acceptable one-factor model (tab. 1). In particular, RMSEA was <0.001 , GFI was 1.000, IFI was 1.000, NFI was 1.000, and CFI was 1.000 (tab. 1). Moreover, standardized regression weights for the four items ranged from 0.77 to 0.94.

Concurrent validity of the Greek version of the S-TIAS was very good since we found statistically significant correlation between the S-TIAS and the AIAS-4 ($r=0.608$, $p<0.001$).

DISCUSSION

We translated and validated the S-TIAS into Greek to

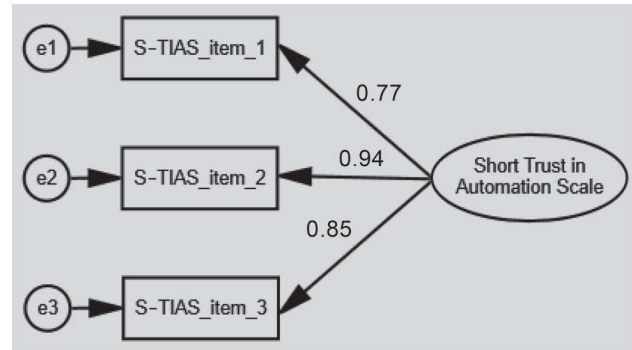


Figure 1. Confirmatory factor analysis for the Greek version of the Short Trust in Automation Scale (S-TIAS).

provide researchers and healthcare professionals with a reliable instrument for assessing trust in artificial intelligence within the general population. The findings indicate that the Greek version of the S-TIAS demonstrates strong psychometric properties, confirming its validity and reliability as a measure of trust toward AI. Furthermore, the analysis supports a unidimensional structure for evaluating trust in AI among the general population.

The Greek version of the S-TIAS demonstrated excellent reliability, with an intraclass correlation coefficient of 0.962 for the test-retest analysis and a Cronbach's alpha of 0.912, indicating strong internal consistency. Confirmatory factor analysis supported a one-factor structure consistent with the original scale, and all fit indices confirmed an excellent model: RMSEA <0.001 , GFI=1.000, IFI=1.000, NFI=1.000, and CFI=1.000. Concurrent validity was also robust, as evidenced by a statistically significant correlation between the S-TIAS and the AIAS-4 ($r=0.608$, $p<0.001$).

Our study faced certain limitations. We relied on a convenience sample from the general population to validate the Greek version of the S-TIAS, which restricts the broader applicability of our results. Future research should aim to validate this instrument using diverse samples within Greece. Moreover, we employed self-reported measures, such as the AIAS-4, to assess the concurrent validity of the S-TIAS. Additional studies could investigate other forms of validity for this tool.

In conclusion, the Greek adaptation of the S-TIAS ex-

Table 1. Confirmatory factor analysis for the Greek version of the Short Trust in Automation Scale.

Model	χ^2	df	χ^2/df	RMSEA	GFI	AGFI	TLI	IFI	NFI	CFI
Three items	0.000	0	NC	<0.001	1.000	NC	NC	1.000	1.000	1.000

NC: Non computable

hibited strong psychometric qualities, confirming its validity and reliability for measuring trust in AI among the general population. By introducing a culturally validated instrument, this study enables systematic monitoring of public and professional attitudes toward AI, supports the assessment of educational and implementation initiatives,

and promotes cross-cultural comparisons with research using the original S-TIAS or other validated tools. Consequently, policymakers and organizational leaders can base AI adoption strategies on robust, locally relevant evidence regarding perceptions of AI's benefits, risks, and appropriate applications.

ΠΕΡΙΛΗΨΗ

Κλίμακα "Short Trust in Automation Scale": Μετάφραση και στάθμιση στην ελληνική γλώσσα

A. ΚΑΤΣΙΡΟΥΜΠΑ,¹ Ι. ΜΩΥΣΟΓΛΟΥ,² Π. ΓΑΛΛΟΣ,³ Ο. ΓΑΛΑΝΗ,³ Μ. ΤΣΙΑΧΡΗ,¹
Π. ΛΙΑΛΙΟΥ,⁴ Ο. ΚΩΝΣΤΑΝΤΑΚΟΠΟΥΛΟΥ,¹ Κ. ΛΑΜΠΡΑΚΟΠΟΥΛΟΥ,⁵ Π. ΓΑΛΑΝΗΣ¹

¹Εργαστήριο Κλινικής Επιδημιολογίας, Τμήμα Νοσηλευτικής, Εθνικό και Καποδιστριακό Πανεπιστήμιο Αθηνών, Αθήνα, ²Τμήμα Νοσηλευτικής, Πανεπιστήμιο Θεσσαλίας, Λάρισα, ³Τμήμα Νοσηλευτικής, Πανεπιστήμιο Δυτικής Αττικής, Αθήνα, ⁴Εργαστήριο Υπολογιστικής Βιοϊατρικής, Τμήμα Ψηφιακών Συστημάτων, Πανεπιστήμιο Πειραιά, Πειραιάς, ⁵Εργαστήριο Σχεδιασμού Ενσωματωμένων Συστημάτων και Εφαρμογών, Τμήμα Ηλεκτρολόγων Μηχανικών και Μηχανικών Υπολογιστών, Πανεπιστήμιο Πελοποννήσου, Πάτρα

Αρχεία Ελληνικής Ιατρικής 2026, 43(6):787–791

ΣΚΟΠΟΣ Η μετάφραση και η στάθμιση της κλίμακας "Short Trust in Automation Scale" (S-TIAS) στα Ελληνικά. **ΥΛΙΚΟ-ΜΕΘΟΔΟΣ** Για τη μετάφραση της κλίμακας στα Ελληνικά εφαρμόστηκε η διαδικασία της «προς τα εμπρός-προς τα πίσω» μετάφρασης. Για τον έλεγχο της αξιοπιστίας της S-TIAS υπολογίστηκε ο συντελεστής Cronbach's alpha. Επί πλέον, εφαρμόστηκε η μέθοδος ελέγχου-επανελέγχου για τον έλεγχο της αξιοπιστίας της S-TIAS, υπολογίζοντας τον συντελεστή ενδοταξικής συσχέτισης. Για τον έλεγχο της εγκυρότητας κατασκευής της S-TIAS πραγματοποιήθηκε η επιβεβαιωτική ανάλυση παραγόντων. Για τον έλεγχο της συγκλίνουσας εγκυρότητας της S-TIAS χρησιμοποιήθηκε η κλίμακα Artificial Intelligence Attitude Scale (AIAS-4). **ΑΠΟΤΕΛΕΣΜΑΤΑ** Διαπιστώθηκε ότι η S-TIAS έχει πολύ καλή αξιοπιστία, καθώς ο συντελεστής ενδοταξικής συσχέτισης στη μέθοδο ελέγχου-επανελέγχου ήταν 0,962 (95% διάστημα εμπιστοσύνης: 0,915–0,983, $p < 0,001$). Επί πλέον, ο συντελεστής Cronbach's alpha για την S-TIAS ήταν 0,912. Η ελληνική έκδοση της S-TIAS είχε έναν παράγοντα, όπως ακριβώς και η πρωτότυπη έκδοση της κλίμακας. Όλοι οι δείκτες επιβεβαίωσαν τη μονοδιάστατη μορφή της S-TIAS. Πιο συγκεκριμένα, ο δείκτης RMSEA ήταν $< 0,001$, ο δείκτης GFI ήταν 1, ο δείκτης IFI ήταν 1, ο δείκτης NFI ήταν 1 και ο δείκτης CFI ήταν 1. Η συγκλίνουσα εγκυρότητα της S-TIAS ήταν πολύ καλή, καθώς βρέθηκε στατιστικά σημαντική συσχέτιση μεταξύ της S-TIAS και της AIAS-4 ($r = 0,608$, $p < 0,001$). **ΣΥΜΠΕΡΑΣΜΑΤΑ** Η κλίμακα S-TIAS είναι ένα αξιόπιστο και έγκυρο εργαλείο για την εκτίμηση της εμπιστοσύνης του γενικού πληθυσμού στην Ελλάδα απέναντι στην τεχνητή νοημοσύνη.

Λέξεις ευρετηρίου: Αξιοπιστία, Εγκυρότητα, Ελληνικά, Κλίμακα "Short Trust in Automation Scale", Τεχνητή νοημοσύνη

References

1. AFROOGH S, AKBARI A, MALONE E, KARGAR M, ALAMBEIGI H. Trust in AI: Progress, challenges, and future directions. *Humanit Soc Sci Commun* 2024, 11:1568
2. OMRANI N, RIVIECCIO G, FIORE U, SCHIAVONE F, AGREDA SG. To trust or not to trust? An assessment of trust in AI-based systems: Concerns, ethics and contexts. *Technol Forecast Soc Change* 2022, 181:121763
3. AQUILINO L, DI DIO C, MANZI F, MASSARO D, BISCONTI P, MARCHETTI A. Decoding trust in artificial intelligence: A systematic review of quantitative measures and related variables. *Informatics* 2025, 12:70
4. BLANCO S. Human trust in AI: A relationship beyond reliance. *AI Ethics* 2025, 5:4167–4180
5. CHIOU EK, LEE JD. Trusting automation: Designing for responsibility and resilience. *Hum Factors* 2023, 65:137–165
6. SCHAEFER KE, CHEN JYC, SZALMA JL, HANCOCK PA. A meta-analysis of factors influencing the development of trust in automation: Implications for understanding autonomy in future systems. *Hum Factors* 2016, 58:377–400
7. HANCOCK PA, BILLINGS DR, SCHAEFER KE, CHEN JYC, DE VISSER EJ, PARASURAMAN R. A meta-analysis of factors affecting trust in human-robot interaction. *Hum Factors* 2011, 53:517–527

8. DEVISSER EJ, PAK R, SHAWTH. From “automation” to “autonomy”: The importance of trust repair in human-machine interaction. *Ergonomics* 2018, 61:1409–1427
 9. GLIKSON E, WOOLLEY AW. Human trust in artificial intelligence: Review of empirical research. *Academy of Management (AOM)* 2020, 14:627–660
 10. JIAN JY, BISANTZ AM, DRURY CG. Foundations for an empirically determined scale of trust in automated systems. *Int J Cogn Ergonomics* 2000, 4:53–71
 11. McGRATH MJ, LACK O, TISCH J, DUENSER A. Measuring trust in artificial intelligence: Validation of an established scale and its short form. *Front Artif Intell* 2025, 8:1582880
 12. GALANIS P. Translation and cross-cultural adaptation methodology for questionnaires in languages other than Greek. *Arch Hellen Med* 2019, 36:124–135
 13. GALANIS P. Validity and reliability of questionnaires in epidemiological studies. *Arch Hellen Med* 2013, 30:97–110
 14. GRASSINI S. Development and validation of the AI Attitude Scale (AIAS-4): A brief measure of general attitude toward artificial intelligence. *Front Psychol* 2023, 14:1191628
 15. WORLD MEDICAL ASSOCIATION. World Medical Association Declaration of Helsinki: Ethical principles for medical research involving human subjects. *JAMA* 2013, 310:2191–2194
 16. BAUMGARTNER H, HOMBURG C. Applications of structural equation modeling in marketing and consumer research: A review. *Int J Res Mark* 1996, 13:139–161
 17. HU LT, BENTLER PM. Fit indices in covariance structure modeling: Sensitivity to underparameterized model misspecification. *Psychol Methods* 1998, 3:424–453
- Corresponding author:*
- P. Galanis, 123 Papadiamantopoulou street, 115 27 Athens, Greece
e-mail: pegalan@nurs.uoa.gr