

한국판 인공지능에 대한 태도 척도(ATTARI-12)의 타당화: 긍·부정 문항 방법효과에 따른 쌍요인 모형 검증

이 선 영¹⁾ 김 효 미[†] 안 현 의²⁾

본 연구에서는 인공지능에 대한 태도를 인지적, 정서적, 행동적 측면에서 포괄적으로 측정하고자 Stein 등(2024)이 개발한 인공지능에 대한 태도 척도(ATTARI-12)를 국내 성인을 대상으로 타당화하였다. ATTARI-12는 세 가지 심리적 구성 요소를 반영하여 문항이 개발되었으나, 척도의 구조는 단일 요인을 이루도록 설계되었다. 본 연구에서는 원 척도를 한국어로 번안한 후, 전국의 만 19~69세 성인 남녀 223명을 대상으로 문항 분석 및 탐색적 요인분석을 실시하였다. 그 결과, 긍정적 표현 문항과 부정적 표현 문항이 각각 독립된 요인으로 나뉘는 2요인 구조로 나타났다. 이를 검증하기 위해 별도의 독립된 223명을 대상으로 확인적 요인분석(CFA)을 실시하였다. 분석에서는 탐색적 요인분석 결과에 기반한 2요인 모형, 부정적 표현 문항의 방법효과를 반영한 쌍요인 S-1 모형, 그리고 수정지수를 사용한 수정 쌍요인 S-1 모형 간의 경쟁 모형 비교를 수행하였다. 그 결과, 구조적으로는 원 척도와 동일한 형태인 쌍요인 S-1 모형이 한국어판에서도 가장 우수한 적합도를 보였으며, 수정지수를 반영한 분석을 통해 이를 더욱 정교화한 수정 쌍요인 S-1 모형을 최종 측정 모형으로 확정하였다. 이후 수렴 타당도, 변별 타당도를 검증한 결과 모두 양호한 것으로 나타났다. 마지막으로 본 연구 결과의 의의와 시사점, 제한점 및 후속 연구에 대한 제언을 논의하였다.

주요어 : 한국판 인공지능에 대한 태도 척도, 탐색적 요인분석, 확인적 요인분석

1) 제1저자 : 이화여자대학교 심리학과 박사 수료

† 교신저자 : 김효미, 이화여자대학교 심리학과 박사 수료, (03760) 서울시 서대문구 이화여대길 32

Tel: 02-3277-2638, E-mail: hyomi0818@daum.net

2) 공동저자 : 이화여자대학교 심리학과 교수



Copyright ©2025, The Korean Psychological Association of Culture and Social Issues
This is an Open Access article distributed under the terms of the Creative Commons Attribution Non-Commercial License
(<http://creativecommons.org/licenses/by-nc/4.0/>) which permits unrestricted non-commercial use, distribution, and reproduction in any medium, provided the original work is properly cited.

날이 갈수록 인공지능 기술이 급속히 발전하며 사람들의 일상생활에 깊숙이 자리 잡음에 따라, 인공지능은 다양한 연구 분야에서 핵심 과제로 주목받고 있다(이장주, 2022). 인공지능(Artificial Intelligence, AI)이란 인간의 지능이 필요한 작업을 자율적으로 수행하며, 학습 및 적응 능력을 갖추고 대량의 데이터를 처리할 수 있는 기술을 의미한다(Stein et al., 2024). 최근에는 인공지능과의 상호작용이 인간 대 인간 상호작용과 유사한 방식으로 더욱 친밀하고 구체적으로 발전하고 있기 때문에(Sundar, 2020; Westerman et al., 2020), 인공지능 기술이 사용자에게 미치는 영향과 사용자의 반응을 심층적으로 이해하는 것이 매우 중요하다.

태도는 행동의 강력한 예측 요인이기 때문에(Ajzen et al., 2018), 인공지능에 대한 태도 연구는 사용자를 이해하고 효과적인 인공지능 관련 연구를 해 나가는 데 필수적이다. 따라서, 인공지능 관련 태도를 평가할 때 사용자가 개념을 명확하게 이해하고 있는지를 확인할 수 있는 적절한 도구의 필요성이 강조된다.

선행 연구에 따르면, 인공지능에 대한 태도와 인식은 혼재되어 있다(Kolasinska et al., 2019; Fast & Horvitz, 2017). 인공지능 기술에 대한 긍정적인 인식은 소비자가 인공지능 제품을 실제로 사용할 때뿐만 아니라, 기술 전반에 대한 태도를 조사하는 일반적인 상황에서도 나타난다. 특히 소비자들은 인공지능 서비스를 접할 때 놀라움과 흥미 등 긍정적인 태도를 가지는 경향이 있으며, 이는 종종 특정 업무나 의사결정을 지원할 수 있는 인공지능의 잠재력에서 비롯된다(Shank et al., 2019; Gerlich, 2023). 반면, 일터에서는 인공지능이

직업 안정성이나 작업자의 통제감과 자율성에 위협을 준다는 인식으로 인해 불안과 같은 부정적 태도가 두드러진다(Yam et al., 2023; Gerlich, 2023; Stein et al., 2024; Nam, 2019). 또한, 국가마다 인공지능 기술에 대한 태도와 인식이 상이한 것으로 나타났다(de la Porte et al., 2024; 정세정, 신영규, 2024). 특히, 한국은 컴퓨터와 생성형 인공지능 사용 비율이 OECD 10개국 중에 가장 높고, 인공지능 소프트웨어 및 장비 사용율은 미국에 이어 2위에 위치하며, 디지털 전환으로 인한 기회를 위협보다 더 높게 인식하는 국가로 분류되지만, 한국인들은 자신들의 디지털 기술에 대한 역량을 낮다고 평가하며, 신기술이 자신의 업무 능력을 위협할 가능성이 크다고 우려하는 경향을 보였다(정세정, 신영규, 2024). 이러한 결과는 인공지능에 대한 태도와 인식이 복잡하고 다차원적임을 보여준다.

해외에서는 인공지능에 대한 태도를 측정하는 다양한 척도가 개발되었으며, 이들은 인공지능과의 상호작용에 대한 사람들의 반응을 포괄적으로 평가하려는 시도를 하고 있다. 대표적으로 인공지능에 대한 태도 척도(GAAIS)(Schepman & Rodway, 2020)는 감정적 반응, 사회적 및 개인적 유용성 평가, 우려 등을 포함하고 있으며, 그 외에도 인공지능에 대한 태도 척도(ATAI)(Sindermann et al., 2021), 인공지능 태도 척도(AIAI)(Krägeloh et al., 2024), 인공지능 불안 척도(AIAS)(Wang & Wang, 2022), 인공지능의 위협 척도(TAI)(Kieslich et al., 2021), 자율 기술에 대한 우려 설문지(Stein et al., 2019), 인공지능에 대한 냉소적 적대감 척도(CHTAIS)(Bochniarz et al., 2022) 등이 개발되었다. 국내에서도 Park 등(2024)이 미국 직장인을 대상으로 개발한 인공지능 기술에 대한 다차

원 척도(AAAW)를 토대로 한국형 인공지능 기술 적용에 대한 직장인 태도 척도(K-AAAW) (김시내, 박지영, 2024)가 개발되었고, Schepman와 Rodway(2020)이 근로자를 대상으로 개발한 인공지능에 대한 일반 태도 척도(GAAIS)를 간호대학생을 대상으로 타당화한 GAAIS-K(서연희, 안정원, 2022) 등이 제시되었다.

이처럼 최근 들어 인공지능에 대한 사람들의 태도를 평가하는 척도들이 꾸준히 개발되고 있으나, 다음과 같은 세 가지 한계점이 있다. 먼저, 심리학적 관점에서 인지, 정서, 행동의 다차원적 요소를 포괄하는 척도를 찾아보기 어렵다. 인공지능에 대한 태도를 측정하는 척도로는 인공지능에 대한 냉소적 적대감 척도(CHTAIS)와 인공지능에 대한 태도 척도(ATAI)가 있으나, 부정적인 태도에 치중하고 있어서 인공지능에 대해서 통합적인 태도를 측정하기에는 한계가 있다. CHTAIS는 인공지능에 대한 냉소적 적대감을 측정하고 있는 도구로 인공지능이 부정적 의도를 가지고 있거나, 신뢰하기 어려운 정도를 평가한다(Bochniarz et al., 2022). ATAI도 인공지능에 대한 두려움, 일자리 손실 등에 초점을 두고 있으며, 총 5가지 문항 중에 3가지 문항이 두려움에 초점을 두고 있다(Sindermann et al., 2021). 인공지능에 대한 태도를 측정하는 척도가 부정적 측면을 위주로 개발하는 한계를 보완하고자 긍정적, 부정적 태도를 모두 포함하는 척도가 개발되기도 했다. 이러한 시도의 일환으로 개발된 AIAI는 총 16문항 중 긍정 문항 8문항, 부정 문항 8문항으로 구성되어, 인공지능에 대한 다면적 인식을 포착하고자 하였다(Krägeloh et al., 2024). 이 척도는 인공지능에 대한 윤리적 측면과 사회적 영향에 관한 문항도 포함하고 있으나, 인지적, 정서적, 행동적 차원

을 아우르는 체계적인 태도 평가에는 한계가 있다.

둘째, 연구의 초점이 특정 집단에 맞추어져 있어 해당 척도가 일반적으로 사용되기에는 제한적이다. 인공지능에 대한 태도 척도 연구는 대부분 중고등학생(Bochniarz et al., 2022)이나 직장인(Schepman & Rodway, 2020) 등 특정 연령대나 특정 집단을 대상으로 진행되고 있다. 이는 각 집단의 특성과 맥락에 따라 척도의 설계와 타당화가 다르게 이루어지다 보니 척도를 전반적인 개인의 인공지능에 대한 태도를 측정하는 데 활용하기는 어렵다. 또한, 특정 연령대나 집단을 기반으로 타당화된 척도가 다른 연령대나 집단에서 동일한 신뢰도와 타당성을 가지는지에 대한 검증이 부족한 경우가 많다. 이는 척도의 일반화 가능성을 제한하며, 다양한 연령대와 배경을 포괄할 수 있는 척도의 개발이 필요하다.

셋째, 인공지능에 대한 태도 척도는 문항 수가 적은 간단한 형태로 개발되는 경우가 많다. 간단한 척도는 빠른 평가와 대규모 연구에서 시간과 자원을 효율적으로 활용할 수 있다는 장점이 있지만, 인공지능에 대한 태도를 면밀하게 다루기는 어렵다. 특히, 인공지능이 사회 및 개인에 미치는 영향을 깊이 있게 분석하거나 태도의 다양한 측면을 포괄적으로 측정하는 데 어려움을 겪을 수 있다. 예를 들어, 인공지능 태도 척도(AIAS-4)는 4개의 문항으로 구성된 간단한 척도로, 인공지능 기술이 개인의 삶과 직업에 미치는 긍정적 영향을 포함하여 사회적 위험과 사용 의도를 평가한다(Grassini, 2023). 이 척도는 단일 요인 구조로 설계되어 빠르고 간단하게 인공지능에 대한 일반적인 태도를 측정하는 데 적합하지만, 보다 복잡한 주제나 세부적인 태도 차원을 다루

기에는 제한적이다. 또 다른 예로, ATAI는 5개의 문항으로 구성되어 있으며, 인공지능에 대한 수용과 두려움을 측정한다(Sindermann et al., 2021). 이 척도 역시 간단한 구조로 설계되어 다양한 문화권에서 빠른 평가를 가능하게 하지만, 인공지능에 대한 복잡한 감정적, 윤리적, 그리고 기술적 측면을 포괄적으로 다루기에는 부족한 점이 있다. 이처럼 문항 수가 적은 척도는 간단한 평가와 대규모 조사에서 유용한 도구로 작용할 수 있으나, 인공지능에 대한 태도를 보다 면밀히 분석하기 위해서는 추가적인 문항 개발과 다양한 요소를 고려한 척도가 필요하다.

이에 본 연구에서는 기존 척도의 한계를 보완하여 국내에서 인공지능에 대한 태도를 보다 신뢰롭고 타당하게 측정할 수 있는 통합적이고 심층적인 도구를 마련하고자 한다. ATTARI-12는 인공지능에 대한 태도를 포괄적으로 측정하기 위해 개발된 도구로, 인지(cognitive), 정서(emotional), 행동(behavioral) 측면의 세 가지 심리적 구성 요소를 균형 있게 반영하여 문항이 구성되었다. 각 영역은 긍정적 태도와 부정적 태도를 동일한 수의 문항으로 측정하도록 설계되었으나, 광범위한 학제 간 적용을 위해 단일 요인 구조를 목표로 개발되었다(Stein et al., 2024). 총점이 높을수록 인공지능에 대해 적극적이고 긍정적인 관심을 나타내며, 점수가 낮을수록 인공지능을 거부하거나 회의적인 태도를 나타낸다. 즉, ATTARI-12의 문항 구성은 인공지능 기술이 개인과 사회에 미치는 영향을 이해하는 인지적 신념, 인공지능과 관련된 감정적 반응, 그리고 인공지능 기술을 사용할 의도 및 관련 행동을 포함하며, 이를 통해 태도의 여러 요소를 포괄적으로 평가할 수 있는 도구로 개발

되었다.

또한 ATTARI-12는 특정 집단에 국한되지 않고 일반 성인을 대상으로 인공지능에 대한 태도를 측정하였다(Stein et al., 2024). 기존 연구들은 중학생, 고등학생, 직장인 등 특정 집단의 맥락에서 인공지능에 대한 태도를 파악하는데 유용했지만, 이를 다른 집단으로 일반화시키기에는 한계가 있었다. 이러한 한계를 보완하고자 ATTARI-12는 특정 집단이나 연령층이 아닌 개인 일반을 대상으로 함으로써, 인공지능에 대한 보편적 태도 이해를 촉진할 수 있다는 장점이 있다. 이러한 척도의 장점을 국내에서도 효과적으로 활용하기 위해서는 한국 문화를 반영하여 적절한 언어적 의미와 표현을 적용한 번안 및 타당화 과정이 필요하다.

따라서 본 연구에서는 ATTARI-12 척도를 한국적 맥락에 맞게 번안하고 타당화하는 과정을 통해 국내에서의 적용 가능성을 검토하였다. 구체적으로, 타문화 심리측정 도구 번안 시 권장되는 역번역 절차(김아영, 임은영, 2003)를 거쳐 문항의 언어적·문화적 적합성을 확보하였으며, 이후 전국의 만 19세에서 만 69세까지의 성인 남녀를 대상으로 온라인 설문조사를 통해 자료를 수집하였다.

연구는 두 단계로 구성되었다. 먼저, 연구 1에서는 한국 문화적 맥락을 고려하여 번안된 ATTARI-12 문항에 대한 문항 분석과 탐색적 요인분석(EFA)을 통해 척도의 요인 구조를 탐색하였다. 이를 통해 ATTARI-12의 문항이 국내 성인 집단에서도 타당하게 구성되는지를 검토하였다. 다음으로, 연구 2에서는 연구 1과는 독립된 표본을 대상으로 확인적 요인분석(CFA)을 실시하여 탐색적 요인분석에서 도출된 구조와 Stein 등(2024)에서 제시된 모형들에

대해 적합성을 검증하였다.

더불어, ATTARI-12 척도의 수렴 타당도를 확인하기 위해 관련성이 높은 변인들과의 상관관계를 분석하였다. 구체적으로, 인공지능에 대한 태도를 측정하는 기존 척도인 인공지능에 대한 태도(GAAIS-K)와의 상관관계를 통해 개념적 유사성을 검토하였다. 두 척도는 모두 인공지능 전반에 대한 긍정적 및 부정적 반응을 평가하기 때문에 높은 정적 상관이 예상된다. 또한, 생성 인공지능 리터러시(K-GAIL), 인공지능 이용 의도와의 상관도 분석하여, 인공지능 태도가 실제 활용 역량과 행동 의도와 어떻게 관련되는지를 확인하고자 하였다. K-GAIL은 인공지능 활용 능력과 윤리적 태도를 포함하고 있으며, 이용 의도는 태도의 행동적 귀결을 설명하는 지표로 기능한다. 이러한 분석을 통해 ATTARI-12가 인접 개념들과 개념적 연관성을 갖는지를 검토하였다.

한편, 변별 타당도를 검증하기 위해 ATTARI-12와 이론적으로 구분되는 구인인 컴퓨터 자기효능감(CSEM)과 인공지능 불안(AIAS)과의 상관관계를 살펴보았다. 컴퓨터 자기효능감은 기술에 대한 전반적인 자기효능감을 측정하는 척도로, 자기효능감과 인공지능 수용 태도 간의 일관된 경향성이 확보되지 않은 것으로 보고된 바 있다(박철우, 강경란, 2022). 또한, 인공지능 불안은 기존 연구(Krägeloh et al., 2024)에서도 인공지능에 대한 태도와 인공지능 불안은 서로 다른 구인으로 낮은 상관을 보이는 것으로 보고되었다. 따라서 이들 변인과의 낮은 상관은 ATTARI-12의 변별 타당도를 지지하는 근거가 된다.

연구 1

연구 1에서는 ATTARI-12를 한국어로 번안하여 예비 문항을 구성하고, 예비 문항을 대상으로 SPSS 23과 Mplus 7.4를 사용하여 문항 분석 및 탐색적 요인분석을 실시했다.

방 법

연구 대상

본 연구는 이화여자대학교 생명윤리위원회의 승인(ewha-202503-0001-01)을 받은 후 진행되었다. 전국에 거주하며 인공지능 기술에 대해 알고 있는 만 19세~만 69세 성인 남녀를 대상으로 온라인 설문조사를 진행하였으며, 리서치 회사에 의뢰해 수집하였다. 총 446명이 설문에 참여했으며, 무작위 추출법을 통해 446명의 절반인 223명의 자료를 연구 1 분석에 사용하였다.

연구 1 참여자의 구성을 보면 성별은 남성이 79명(35.4%), 여성이 144명(64.6%)이었다. 연령은 만30-39세가 69명(30.9%), 만40-49세가 57명(25.6%), 만19-29세가 49명(22%), 만50-59세가 42명(18.8%), 만60-69세가 6명(2.7%) 순이었다. 최종 학력은 대졸 138명(61.9%), 전문대졸 30명(13.5%), 고졸 28명(12.6%), 대학원 이상 27명(12.1%) 순이었다. 현재 직업은 사무직이 91명(40.8%), 전문직 26명(11.7%), 서비스, 판매직 21명(9.4%), 전업주부와 학생이 각각 18명씩(8.1%)을 차지했다. 정기적으로 사용하는 기기, 기술(중복 선택 가능)은 휴대용 기기 > 컴퓨터 > 생성형 AI > 인공지능을 사용하는 소프트웨어, 장비 > 물리적 작업을 수행하기 위한

자동화 장비 순이었다. 인공지능 관련 교육을 받은 경험은 ‘있음’이 67명(30%), ‘없음’ 156명(70%)였다. ‘자신이 가진 인공지능에 대한 지식수준에 대해 어느 정도 알고 있다’는 181명(81.2%), ‘충분히 알고 있다’는 22명(9.9%), ‘전혀 모른다’는 17명(7.6%), ‘자세히 알고 있다’는 3명(1.3%) 순이었다.

측정 도구

인공지능에 대한 태도(Attitudes towards artificial intelligence Scale, ATTARI-12)

Stein 등(2024)이 개발한 ATTARI-12 척도를 사용하였다. 이 척도는 총 12개 문항으로 구성되어 있으며, 인지적, 정서적, 행동적 태도의 3가지 심리적 측면마다 각각 긍정적 표현 문항 2개와 부정적 표현 문항 2개로 구성되어 있다. Likert 5점 척도로 평가되고 부정적 표현의 경우 역채점이 필요하며, 총 합계 점수가 높을수록 인공지능에 대한 태도가 긍정적임을, 점수가 낮을수록 부정적임을 의미한다. Stein 등(2024)의 연구에서 내적 합치도(Cronbach's α)는 .93 이었으며, 본 연구에서 내적 합치도(Cronbach's α)는 .89이었다.

문항 번안

ATTARI-12 원 척도 사용에 대해 허가를 받기 위해 1저자인 Jan-Philipp Stein에게 연락하여 척도 사용에 대한 허가를 받았다. 이후 본 연구 저자 2명이 영어-한국어 1차 번역을 하고, 영어권 국가에서 심리학 관련 전공의 석사 졸업생이자 한국에서 심리학 전공으로 박사 졸업생 1인에게 검토받았다. 이후 전문 번역가가 다시 역번역을 하였고, 번역본과 역번

역본을 연구자가 비교함으로써 불일치하는 것으로 보이는 문항들은 같이 1차 검토한 심리학 박사 졸업생 1인과 함께 수정하였다. 재차 수정 과정을 거친 번역본과 원 척도, 그리고 역번역본을 토대로 미국에서 장기 거주 경험이 있으며 영어에 능숙한 심리학 전공 교수 1인의 검토를 거쳐, 보다 명확하면서도 국내에서 친숙하게 받아들여질 수 있는 표현으로 문장을 수정하여 최종 문항을 확정하였다. ATTARI-12 원 척도 문항과 번역본을 비교하여 원 척도의 문장 표현을 변질시키지 않으면서 사람들이 쉽게 이해할 수 있도록 문화적 배경을 고려하여 예비 척도 문항을 확정하였다.

결 과

문항의 양호도 검증

예비 척도의 12개 문항의 양호도를 검증하고자 각 문항의 평균과 표준편차, 문항 간 상관, 문항-총점 간 상관 및 내적 합치도(Cronbach's α)를 확인하였다. 본 연구에서 사용된 척도는 5점 Likert 척도로 이루어져 있으므로 문항 평균이 1.0 미만이거나 4.0을 초과하는 경우와 문항의 표준편차가 .7 이하인 경우, 문항 간 상관이 .2 이하이거나 .8 이상인 경우(김남걸, 2001), 문항-총점 간 상관이 .4 미만인 경우(Ware & Gandek, 1998) 문항의 양호도가 떨어진다고 판단하여 문항의 제외를 고려하였다. 또한 내적 합치도 분석에서 문항 제거 시 Cronbach's α 값이 상승하는 경우, 해당 문항이 척도의 전반적인 신뢰도를 저해한다고 보고 제외를 고려하였다.

추가로, 표본 자료의 정규성 및 문항의 변별도를 확인하기 위해 각 문항의 최소 및 최대값, 왜도 및 첨도를 검토하였다. 단변량 정규성 평가에서 Curran 등(1996)이 제안한 기준에 따라 왜도의 절댓값이 2를 초과하거나 첨도의 절댓값이 7을 초과하면 정규성 가정을 위배하는 것으로 판단하였다. 본 연구의 모든 문항에서 왜도의 절댓값은 .23~.92, 첨도의 절댓값은 .09~2.18의 범위로 나타나, 정규성 가정을 충분히 충족하였다.

문항의 변별도를 확인하기 위해 최소 및 최대값을 살펴본 결과, 모든 문항의 최소값은 1점, 최대값은 5점으로 나타나 5점 Likert 척도의 전 범위를 고르게 사용한 것으로 나타났다. 문항의 최소-최대값 범위는 4로, 지나치게 좁은 범위(3 이하)를 보이는 문항은 없었다. 평균과 표준편차를 확인한 결과, 문항 2번(“나는 AI에 의존하는 기술을 사용하고 싶다”)의 평균이 4.14로 4.0을 초과하고, 표준편차는 .70으로 기준(.7 이하)에 근접하여 충족하지 못하였다. 그러나 문항-총점 간 상관이 .48으로 적절한 수준이었고, 문항 제거 시 Cronbach's α 값이 낮아져 문항을 유지하기로 하였다. 그 외의 문항들은 평균, 표준편차, 문항 간 상관, 문항-총점 간 상관 등 모든 기준을 충족하였다. 따라서 본 연구에서는 모든 문항을 유지하고 탐색적 요인분석을 실시하였다.

탐색적 요인분석

SPSS 23.0을 활용하여 변안된 12문항이 요인분석에 적합한지 확인한 결과, KMO 표본 적합도(Kaiser-Meyer-Olkin Measure of Sampling Adequacy)값이 .904, Bartlett의 구형성 검정(χ^2

=1152.082, $df=66$, $p<.001$)에서도 적합한 것으로 나타나 수집된 자료는 요인분석에 적합한 것으로 확인하였다. 이후 탐색적 요인분석은 보다 정확한 요인구조의 확인을 위해 Mplus 7.4 프로그램을 이용하여 실시하였으며, 최대우도법(Maximum Likelihood, ML)과 요인의 상관을 허용하는 사각회전법(Oblimin)을 적용하여 수행하였다. 모형의 적합도는 Hu와 Bentler(1999)의 기준에 따라 RMSEA는 .08 이하, CFI와 TLI는 .90 이상, SRMR은 .08 이하일 때 적합도가 양호한 것으로 판단하였다.

분석 결과, 단일요인 모델($\chi^2=221.995$, $df=54$, RMSEA=.118, CFI=.850, TLI=.816, SRMR=.070)보다는 2요인 모델($\chi^2=86.797$, $df=43$, RMSEA=.068, CFI=.961, TLI=.940, SRMR=.032)의 적합도가 통계적으로 유의하게 더 양호하였다($\Delta\chi^2=135.198$, $\Delta df=11$, $p<.001$). 3요인 모델은 적합도가 다소 개선되었으나($\chi^2=54.980$, $df=33$, RMSEA=.055, CFI=.980, TLI=.961, SRMR=.025), 요인의 해석 가능성 및 요인 간 상관을 고려할 때, 2요인 모델이 가장 적절한 것으로 판단되었다. 2요인 모델은 전체 변량의 58.25%를 설명하였다. 요인 1의 설명량은 46.67%로, 요인 1에 속한 문항의 부하량은 .578에서 .796 사이로 나타났고, 부정적 태도를 나타낸 문항들로 구성되었다. 요인 2의 설명량은 11.58%로, 요인 2에 속한 문항 부하량은 .436에서 .755 사이로 나타났으며, 긍정적 태도를 나타낸 문항들로 구성되었다. 두 요인 간 상관은 .64로 중간 수준의 정적 상관을 나타내었다. 각 요인에 해당되는 문항은 표 1과 같다.

표 1. 탐색적 요인분석(EFA)을 통한 문항별 요인부하량 및 최종 문항(영어 원문 및 한글 번역문)

	문항(R-역채점)	요인부하량	
		1	2
2R	나는 AI에 대해 강한 거부감(부정적인 감정)을 느낀다. (I have strong negative emotions about AI.)	.626	
4R	AI는 장점보다 단점이 더 많다고 생각한다. (AI has more disadvantages than advantages.)	.728	
7R	나는 AI 없이 작동하는 기술을 더 선호한다. (I prefer technologies that do not feature AI.)	.796	
8R	나는 AI에 대해 두려움을 느낀다. (I am afraid of AI.)	.652	
10R	AI는 문제를 해결하기보다는 새로운 문제를 발생시키는 경우가 많다. (AI creates problems rather than solving them.)	.734	
12R	나는 AI를 기반으로 한 기술을 가급적 사용하지 않으려 한다. (I would rather avoid technologies that are based on AI.)	.578	
1	AI는 세상을 더 나은 곳으로 변화시킬 것이다. (AI will make this world a better place.)		.703
3	나는 AI를 활용한 기술을 사용해 보고 싶다. (I want to use technologies that rely on AI.)		.436
5	나는 앞으로 AI가 어떻게 발전될지 기대된다. (I look forward to future AI developments.)		.557
6	AI는 세상의 다양한 문제에 대한 해결책을 제시한다. (AI offers solutions to many world problems.)		.755
9	나는 AI 없이 사용되는 기술보다 AI가 적용된 기술을 더 선택하고 싶다. (I would rather choose a technology with AI than one without it.)		.728
11	나는 AI에 대해 대체로 긍정적인 감정을 갖고 있다. (When I think about AI, I have mostly positive feelings.)		.657

연구 2

연구 2에서는 연구 1에서 탐색적 요인분석 결과로 도출된 2요인 구조와 Stein 등(2024)이 제시한 다양한 구조 모형의 적합도를 검증하기 위해 연구 1과 독립된 전국의 성인 223명

의 자료를 대상으로 SPSS 23.0과 Mplus 7.4를 사용하여 문항 양호도 분석 및 확인적 요인분석을 실시하였다. 최종 모형을 선택한 후, 기존의 인공지능에 대한 태도 척도(GAAIS-K), 한국형 생성 인공지능 리터러시 척도(K-GAIL)와 인공지능 이용 의도 척도와 의 상관 분석을 통

해 수렴 타당도를, 컴퓨터 자기 효능감 척도(CSEM), 인공지능 불안 척도(AIAS)와 상관 분석을 통해 변별 타당도를 확인하였다.

‘전혀 모른다’는 16명(7.2%), ‘자세히 알고 있다’는 2명(0.9%) 순이었다.

측정 도구

방 법

인공지능에 대한 태도(Attitudes towards artificial intelligence Scale, ATTARI-12)

연구 대상

Stein 등(2024)이 개발한 ATTARI-12 척도를 번역 및 역번역 과정을 통해 번안하여 연구 1과 동일한 척도를 사용하였다. 총 12개 문항으로 구성되어 있으며, 인지적, 정서적, 행동적 태도의 3가지 심리적 측면마다 각각 긍정적 표현 문항 2개와 부정적 표현 문항 2개로 구성되어 있다. Likert 5점 척도로 평가되고 부정적 표현의 경우 역채점이 필요하며, 총 합계 점수가 높을수록 인공지능에 대한 태도가 긍정적임을, 점수가 낮을수록 부정적임을 의미한다. Stein 등(2024)의 연구에서 내적 합치도(Cronbach's α)는 .93 이었으며, 본 연구 1에서 내적 합치도(Cronbach's α)는 .89이었다.

연구 1과 동일하게 전국에 거주하며 인공지능 기술에 대해 알고 있는 만 19세~만 69세 성인 남녀를 대상으로 온라인 설문조사를 진행하였으며, 리서치 회사에 의뢰해 수집하였다. 무작위 추출법을 통해 연구 1에서 사용한 446명의 절반과 다른 223명의 자료를 연구 2 분석에 사용하였다.

인공지능에 대한 태도(Korean version of General Attitudes towards Artificial Intelligence Scale, GAAIS-K)

연구 2 참여자의 구성을 보면 성별은 남성이 80명(35.9%), 여성이 143명(64.1%)이었다. 연령은 만30-39세가 79명(35.4%), 만40-49세가 63명(28.3%), 만19-29세가 40명(17.9%), 만50-59세가 34명(15.2%), 만60-69세가 7명(3.1%) 순이었다. 최종 학력은 대졸 131명(58.7%), 고졸 35명(15.7%), 전문대졸 28명(12.6%), 대학원 이상 29명(13.0%) 순이었다. 현재 직업은 사무직 94명(42.2%), 전문직 28명(12.6%), 전업주부 23명(10.3%), 학생 18명(8.1%), 서비스, 판매직 16명(7.2%)을 차지했다. 정기적으로 사용하는 기기, 기술(중복 선택 가능)은 휴대용 기기 > 컴퓨터 > 생성형 AI > 인공지능을 사용하는 소프트웨어, 장비 > 물리적 작업을 수행하기 위한 자동화 장비 순이었다. 인공지능 관련 교육을 받은 경험은 ‘있음’이 66명(29.6%), ‘없음’ 157명(70.4%)였다. ‘자신이 가진 인공지능에 대한 지식수준에 대해 어느 정도 알고 있다’는 183명(82.1%), ‘충분히 알고 있다’는 22명(9.9%),

GAAIS는 근로자를 대상으로 AI에 대한 일반적인 태도를 측정하기 위해 개발된 척도로, 인공지능에 대한 긍정적인 측면과 부정적인 측면을 모두 측정하는 척도이다. 단순히 이용 의도와 같은 행동적 측면을 넘어서 인공지능에 대한 윤리적, 사회적 영향을 모두 다루는 척도로 볼 수 있다. 긍정적 태도는 인공지능의 사회적 및 개인적 유용성, 긍정적 감정을 의미하며, 부정적 태도는 인공지능의 위험성, 불안감, 윤리적 문제를 포괄하여 다루고 있다. 본 연구에서는 Schepman와 Rodway(2020)의

척도를 번안하고 타당화한 서연희와 안정원(2022)의 GAAIS-K 척도를 사용하였다. 총 20문항이며 긍정적 태도 12문항, 부정적 태도 8개 문항으로 이뤄져 있으며, 5점 Likert 척도로 '전혀 그렇지 않다(1점)'에서 '매우 그렇다(5점)'로 측정하였다. 점수가 높을수록 긍정적 태도 또는 부정적 태도가 높다는 것을 의미한다. 서연희와 안정원(2022)에서 내적 합치도(Cronbach's α)는 긍정적 태도 .86, 부정적 태도 .74로 나타났으며, 본 연구에서 내적 합치도(Cronbach's α)는 긍정적 태도 .87, 부정적 태도 .84로 나타났다.

한국형 생성 인공지능 리터러시 척도 (Korean Generative AI Literacy Scale, K-GAIL)

생성형 인공지능의 활용 능력을 평가하기 위해 노환호 등(2024)이 개발한 한국형 생성 인공지능 리터러시 척도(K-GAIL)를 사용하였다. K-GAIL은 네 가지 요인으로 구성되는데, 각각 AI 활용능력, 비판적 평가, 윤리적 사용, 창의적 활용이다. AI 활용능력(AI Utilization Ability)은 인공지능을 효과적으로 사용하여 필요한 작업을 수행하는 능력을 측정한다. 비판적 평가(Critical Evaluation)는 생성된 정보를 비판적으로 분석하고, 사실과 의견을 구분하여 사회에 미치는 영향을 평가하고, 윤리적 사용(Ethical Use)은 인공지능을 사용할 때 발생할 수 있는 윤리적 문제를 이해하고, 이에 적절히 대응할 수 있는 능력을 측정한다. 마지막으로 창의적 활용(Creative Utilization)은 인공지능을 창의적으로 활용하여 새로운 아이디어를 제시하거나 문제를 해결하는 능력을 평가한다. K-GAIL 점수가 높을수록 생성 인공지능을 이해하고 활용하는 능력뿐만 아니라 그에 대한 비판적 평가, 윤리적 책임 인식, 창의적 활용

역량이 높음을 의미한다. K-GAIL은 총 12개의 문항으로 구성되어 있으며, 모든 문항은 5점 Likert 척도로 '전혀 그렇지 않다(1점)'에서 '매우 그렇다(5점)'로 측정하였다. 노환호 등(2024) 연구에서 내적 합치도(Cronbach's α)는 .75~.89 수준으로 나타났으며, 본 연구에서 내적 합치도(Cronbach's α)는 .86으로 나타났다.

이용 의도 척도(Intention to use scale)

기술 수용 모델을 기반으로 Venkatesh와 Davis(2000)이 개발한 척도를 황영훈 등(2020)이 번안한 이용 의도 문항을 참고하여 본 연구에서 사용하였다. 이용 의도는 특정 시스템이나 기술을 사용하려는 개인의 행동 의지로 정의할 수 있으며 총 3문항으로 측정하였다. 원래 척도의 '시스템(system)'이라는 표현을 본 연구의 맥락에 맞게 'AI'로 변경하여 사용하였으며, 문항은 7점 Likert 척도로 '전혀 그렇지 않다(1점)'에서 '매우 그렇다(7점)'로 측정하였다. 점수가 높을수록 이용 의도가 높다는 것을 의미한다. 황영훈 등(2020) 연구에서 내적 합치도(Cronbach's α)는 .96이었으며, 본 연구에서 내적 합치도(Cronbach's α)는 .93이다.

컴퓨터 자기 효능감(Computer Self-Efficacy Measure, CSEM)

Compeau와 Higgins(1995)이 개발한 컴퓨터 자기효능감(CSEM) 척도는 개인이 컴퓨터 기술을 학습하고 사용하는데 가지는 자신감을 측정하는데, 다양한 응용 분야에서 단순히 컴퓨터 자기효능감을 넘어서 기술에 대한 자기효능감을 측정하는데 사용되고 있다. 본 연구에서는 김보나 등(2010)이 번안하고 타당화한 척도를 사용하여 연구를 진행하였다. CSEM은 문항은 5점 Likert 척도로 '매우 자신없음(1점)'에

서 ‘매우 자신있음(5점)’로 측정하였으며 총 8 문항이다. 점수가 높을수록 기술에 대한 자기 효능감이 높다는 것을 의미하며, 김보나 등(2010) 연구에서 내적 합치도(Cronbach's α)는 .92였다. 본 연구에서 내적 합치도(Cronbach's α)는 .90으로 나타났다.

인공지능 불안 척도(Artificial Intelligence Anxiety Scale, AIAS)

인공지능 불안을 측정하기 위해 Wang와 Wang(2022)이 개발하고, 정민영 등(2023)이 변안한 21문항을 사용한다. 인공지능 불안 측정 척도는 인공지능 학습에 대한 불안 8문항, 인공지능의 직업 대체에 대한 불안 6문항, 사회 기술적 무지에 대한 불안 4문항, 인공지능 형상에 대한 불안 3문항으로 구성되어 있다. 각 문항은 5점 Likert 척도로 측정되었으며 ‘매우 그렇지 않다(1점)’에서 ‘매우 그렇다(5점)’으로 측정하였다. 점수가 높을수록 인공지능 불안을 높음을 의미한다. 박선규 등(2024)에서 문항 전체의 내적 합치도(Cronbach's α)는 .93으로 나타났고, 학습에 대한 불안 .93, 인공지능의 직업 대체에 대한 불안 .89, 사회 기술적 무지에 대한 불안 .83, 인공지능 형상에 대한 불안 .84이다. 본 연구에서 문항 전체의 내적 합치도(Cronbach's α)는 .94이고, 하위요인별로는 학습에 대한 불안 .93, 인공지능의 직업 대체에 대한 불안 .91, 사회 기술적 무지에 대한 불안 .85, 인공지능 형상에 대한 불안 .87으로 나타났다.

결 과

확인적 요인분석

SPSS 23.0을 활용하여 연구 1과 동일한 문항 분석 과정을 거쳤으며, 각 문항의 평균, 표준편차, 문항 간 상관, 문항-총점 간 상관이 모두 양호했다.

신뢰도 분석 결과 Cronbach's α 계수는 전체 .86로 각 문항들이 측정하는 데 문제가 없음을 확인하였다. 추가적으로 정규성 확인을 위해 각 문항의 왜도와 첨도를 검토한 결과, 각 문항의 왜도는 -.95~.03, 첨도는 -.82~3.12으로 나타나 왜도 절대값 2 미만, 첨도 절대값 7 미만이라는 Curran 등(1996)의 기준을 충족하였다. 또한 문항의 변별도를 평가하기 위해 최소 및 최대값을 확인한 결과, 최소값은 1~2 점, 최대값은 모두 5점으로 문항의 응답 범위가 충분히 활용되었음을 확인하였다. 따라서 모든 문항을 유지한 상태에서 Mplus 7.4로 확인적 요인분석을 실시하였다.

구조방정식 모형의 적합도는 TLI와 CFI는 상대적 적합도 지수로서 .90 이상이면 좋은 적합도를 나타내며, RMSEA는 절대적 적합도 지수로서, RMSEA 수치는 .05 보다 작으면 좋은 적합도, .08 보다 작으면 괜찮은 적합도, .10 보다 작으면 보통 적합도, .10 보다 크면 나쁜 적합도로 간주된다(홍세희, 2000). 한국판 인공지능에 대한 태도 척도에 대해 최대우도 방법으로 확인적 요인분석(CFA)을 실시한 결과, 표 2와 같이 나타났다.

탐색적 요인분석을 통해 도출된 2요인 모형은 $\chi^2=121.130(df=53, p<.001)$, CFI=.922, TLI=.902, SRMR=.050, RMSEA=.076[90% CI=.058~.094]로 전반적으로 양호한 적합도를 보였다. 그러나 이 2요인 모형의 경우, 문항들이 긍정적 또는 부정적 방식으로 표현되는 형식 자체

표 2. 인공지능에 대한 태도(ATTRA-12) 모형의 적합도 지수

모형	χ^2	df	TLI	CFI	SRMR	RMSEA (90% 신뢰 구간)
1요인	236.384	54	.744	.790	.078	.123(.107~.139)
2요인(긍정, 부정)	121.130	53	.902	.922	.050	.076(.058~.094)
3요인(정서, 행동, 인지)	203.357	51	.773	.825	.078	.116(.099~.133)
쌍요인 S-1 모형 (부정 문항 특수요인)	109.668	48	.902	.929	.044	.076(.057~.095)
수정 쌍요인 S-1 모형 (부정 문항 특수요인)	74.855	46	.952	.967	.037	.053(.030~.074)

가 별도의 특수한 요인으로 작용하여 문항의 실제적인 내용이나 본질적인 의미와 관계없이 응답자들의 답변 경향성에 영향을 미쳤을 가능성을 배제할 수 없다. 따라서 문항의 표현 방식을 고려하여, 문항의 본질적 의미를 나타내는 일반요인과 함께 부정적 문항 표현에 기인한 특수요인을 분리하여 설명하는 쌍요인 S-1 모형을 추가적으로 검증하였다.

부정적으로 표현된 문항의 영향을 고려한 쌍요인 S-1 모형에서는 $\chi^2=109.668(df=48, p<.001)$, CFI=.929, TLI=.902, SRMR=.044, RMSEA=.076[90% CI=.057~.095]으로 기존 2요인 모형보다 다소 개선된 적합도를 나타냈다. 그러나 더욱 정교한 모형을 도출하고 적합도를 향상시키기 위해 수정지수(modification index)를 참고하여 모형 수정의 타당성과 이론적 근거를 면밀히 검토하였다. 수정지수를 이용한 모형 수정에는 제약이 있지만, 동일 변수 내 측정오차 간 공분산 허용은 가능하므로 수정지수가 가장 높게 나타난 동일 변수 내의 오차 간을 순차적으로 공분산을 허용하였다.

오차 간 공분산 허용은 동일 변수 내 문항들이 표현 방식이나 내용적 측면에서 유사성을 지니고 있을 때 발생할 수 있는 잔여 상관을 효과적으로 통제하고 모형의 정확성을 향상시키기 위한 이론적·방법론적으로 타당한 접근이다. 본 연구에서는 수정지수를 바탕으로 6번 문항과 1번 문항, 7번 문항과 4번 문항의 오차 간을 순차적 공분산을 허용하여 모형을 수정한 결과, 수정 쌍요인 S-1 모형의 적합도는 $\chi^2=74.855(df=46, p=.004)$, CFI=.967, TLI=.952, SRMR=.037, RMSEA=.053[90% CI=.030~.074]으로 나타나 모든 적합도 지표에서 가장 우수한 결과를 보였을 뿐 아니라, 이론적 해석 가능성도 우수하여, 수정 쌍요인 S-1 모형을 최종적으로 채택하였다. 한국판 인공지능에 대한 태도 척도에 대한 경로 모형은 그림 1과 같다.

반면, 원 논문에서 제안한 내용적 하위 요인(contents facets)만 고려하거나 내용적 하위 요인과 부정적 문항(item wording)을 모두 고려한 쌍요인 S-1 모델은 낮은 적합도가 도출되

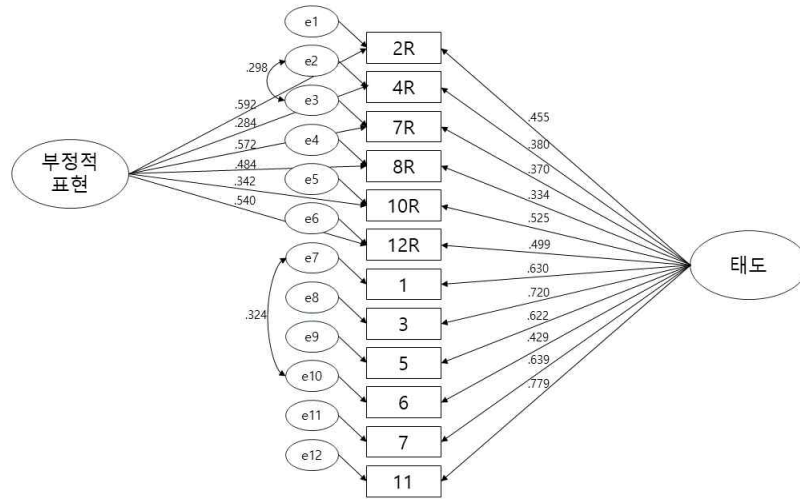


그림 1. 한국판 인공지능에 대한 태도 척도 수정 쌍요인 S-1 모형(표준화 회귀계수)

었다. 이러한 결과는 인지, 정서, 행동 차원의 내용적 하위 요인들이 문항 수준에서 충분히 다뤄지지 않았을 가능성이 있다는 것을 의미한다.

수렴 타당도

ATTARI-12의 수렴 타당도를 확인하기 위해 기존의 인공지능에 대한 태도 척도(GAAIS-K)와의 상관을 분석하였다(표 3). 그 결과, ATTARI-12와 GAAIS-K 전체 점수 간 상관은 매우 높은 정적 상관을 보였다($r=.94$, $p<.01$). GAAIS-K의 하위 요인과의 상관에서도 ATTARI-12는 긍정적 태도 요인과 $.93(p<.01)$, 부정적 태도 요인과 $.52(p<.01)$ 으로 유의한 정적 상관을 나타내었다. 또한, ATTARI-12와 생성 인공지능 리터러시(K-GAIL) 간의 상관을 분석한 결과, 전체 점수 간 유의한 정적 상관이 나타났다($r=.11$, $p<.01$). 하위 요인 중 AI 활용 능력($r=.52$, $p<.01$) 및 창의적 활용($r=.52$, $p<.01$)과는 유의한 정적 상관을 보였다.

며, 윤리적 사용($r=.19$, $p<.05$)도 낮은 수준이지만 유의한 정적 상관을 나타냈다. 반면, 비판적 평가($r=.13$)는 유의하지 않았다. 이는 ATTARI-12가 AI 활용 능력과 창의적 활용과 일정 부분 관련되어 있음을 시사한다. 아울러, 이용 의도와 상관은 매우 높은 정적 상관($r=.74$, $p<.01$)이 나타났으며, 이는 ATTARI-12가 실제 AI 이용 의도와 밀접하게 연관되어 있다는 것을 의미한다. 이러한 결과는 ATTARI-12가 기존 인공지능 관련 태도, 리터러시, 이용 의도와의 정적 상관을 통해 전반적으로 수렴 타당도를 갖추고 있음을 시사한다.

변별 타당도

ATTARI-12의 변별 타당도를 확인하기 위해 컴퓨터 자기효능감(CSEM) 및 인공지능 불안(AIAS) 척도와의 상관을 분석하였다(표 3). 그 결과, ATTARI-12와 CSEM 전체 점수 간에는 유의한 정적 상관이 나타났으나 상관

표 3. 인공지능에 대한 태도(ATTARI-12)와 다른 척도들 간의 상관(N=223)

		ATTARI-12	평균	표준편차
인공지능에 대한 태도 (GAAIS-K)	전체	0.94**	3.58	0.79
	긍정적 태도	0.93**	3.84	0.72
	부정적 태도	0.52**	3.06	0.91
생성 인공지능 리터러시 (K-GAIL)	전체	0.11**	3.54	0.79
	AI활용 능력	0.52**	3.89	0.38
	비판적 평가	0.13	3.35	0.55
	윤리적 사용	0.19*	3.47	0.62
	창의적 활용	0.52**	3.52	0.50
이용 의도	전체	0.74**	5.52	1.00
컴퓨터 자기효능감(CSEM)	전체	0.12**	3.62	0.82
인공지능 불안(AIAS)	전체	-0.30**	3.13	0.93
	인공지능 학습에 대한 불안	-0.33**	2.39	0.89
	인공지능 직업 대체에 대한 불안	-0.22**	3.32	1.04
	사회기술적 무지에 대한 불안	-0.26**	3.44	1.00
	인공지능 형상에 대한 불안	-0.14	3.34	0.95

** $p < .01$, * $p < .05$

계수는 낮은 수준이었다($r = .12$, $p < .05$). 또한 ATTARI-12는 인공지능 불안(AIAS) 전체 점수와 유의한 부적 상관($r = -.30$, $p < .01$)을 보였으며, 하위 요인 중 인공지능 학습($r = -.33$, $p < .01$), 직업 대체($r = -.22$, $p < .01$), 사회기술적 무지($r = -.26$, $p < .01$)에 대한 불안과도 모두 유의한 부적 상관을 나타냈다. 하지만, 모든 상관 계수가 $-.33$ 이하로 낮은 수준에 머물렀으며, 이는 ATTARI-12가 인공지능 불안이나 컴퓨터 자기효능감과 구별되는 태도적 특성을 측정하고 있음을 시사한다. 따라서 ATTARI-12는 관련 개념들과의 낮은 상관을 통해 변별 타당도가 확보되었음을 보여준다.

논 의

본 연구에서는 Stein 등(2024)이 개발한 ATTARI-12를 한국판으로 번안하여 만 19세~만 69세의 성인 남녀를 대상으로 설문을 실시하고, 그 결과를 분석하여 타당화 과정을 수행하였다. 본 연구는 두 가지의 하위 연구로 구성되었다. 먼저 연구 1에서는 원 척도를 한국어로 번안한 후 문항 분석 및 탐색적 요인 분석을 실시하여 문항의 타당성을 평가하였다. 그 결과, 긍정적 태도와 부정적 태도라는 두 가지 하위 요인으로 구성된 총 12문항의 한국판 ATTARI-12 척도가 도출되었다. 연구 2에서는 연구 1에서 제시된 2요인 구조와 Stein 등

(2024)이 제시한 다양한 구조 모형을 바탕으로 확인적 요인분석을 실시하였다. 그 결과, 부정적 문항 표현으로 인해 발생하는 특수요인을 별도로 고려한 쌍요인 S-1 모형이 가장 적합한 것으로 나타났으며, 수정지수를 통해 일부 문항 간 측정 오차의 공분산을 허용한 수정 쌍요인 S-1 모형을 최종 모형으로 선택하였다. 또한, 척도의 타당성을 보다 명확히 검증하기 위해 수렴 타당도와 변별 타당도를 분석하였는데, 이를 위해 기존의 인공지능에 대한 태도(GAIS-K), 한국형 생성 인공지능 리더러시(K-GAIL)와 이용 의도, 컴퓨터 자기효능감(CSEM), 인공지능 불안(AIAS)과 상관 분석을 실시하였다.

본 연구의 주요 결과의 요약과 시사점은 다음과 같다. 첫째, ATTARI-12 원척도를 개발한 Stein 등(2024)의 연구에서는 탐색적 요인분석이 실시되지 않았으나, 본 연구에서는 한국 문화적 맥락을 고려하여 문항 분석과 탐색적 요인분석을 통해 척도의 내적 구조를 탐색하였다. 그 결과 부정적 표현/태도와 긍정적 표현/태도로 구성된 2요인 구조가 도출되었다. Stein 등(2024)은 척도 개발 당시, 부정적 태도에서 긍정적인 태도까지 태도의 전 범위를 편향 없이 측정하기 위해 인지적, 정서적, 행동적 측면에서 각각 긍정적 표현 문항 2개와 부정적 표현 문항 2개로 균형 있게 구성하였다. 그러나 긍정적 표현과 부정적 표현 문항들이 각각 독립된 요인으로 구분된 것은 문항의 표현 방식으로 인해 응답자의 반응에 체계적 편향이 발생하는 방법효과(method effect)로 설명될 수 있다. 기존 선행 연구들에서도 긍정적 표현과 부정적 표현의 혼합 사용으로 인해 응답자들이 체계적인 반응 편향을 나타내는 방법효과가 발생할 수 있으며, 이러한 방법효과

가 척도의 개념적 의도와 다른 요인구조를 유도할 수 있음이 보고되었다(최수미, 조영일, 2013). 또한, 긍정적 및 부정적 진술문이 혼합된 심리적 평정 척도의 경우, 전체 신뢰도가 낮아지거나(Schriesheim & Hill, 1981), 측정오차가 증가할 가능성이 지적되기도 하였다(Tepper & Tepper, 1993). 탐색적 요인분석이 완전한 요인 구조들을 구별할 수 없다는 점을 고려해 봤을 때(Quilty et al., 2006), 본 연구의 연구 2에서는 확인적 요인분석을 통해 긍정 문항과 부정 문항으로 구성된 요인 구조를 포함한 다양한 요인 구조를 체계적으로 평가하고 비교·검증하였다.

둘째, 연구 2에서는 연구 1에서 도출된 2요인 구조(긍정적 태도, 부정적 태도)와 원 개발자인 Stein 등(2024)이 제시한 다양한 모형들을 대상으로 확인적 요인분석(CFA)을 통해 최적의 구조를 검증하였다. 그 결과, 탐색적 요인 분석에서 제안된 2요인 모형은 통계적으로 양호한 수준이었으나, 문항 표현 방식 자체가 별도의 특수요인으로 작용하여 응답에 편향을 초래할 가능성이 존재하였다. 이에 문항 표현 방식으로 인한 방법효과를 명확히 통제하고자, 일반요인과 부정적 문항 표현 특수요인을 동시에 고려한 쌍요인 S-1 모형을 추가적으로 검증하였다. 모형의 적합성을 더욱 개선하고자 수정지수(modification index)를 참고하여 일부 문항의 표현이나 내용적 유사성에서 비롯된 측정 오차 간 공분산을 허용하였다. 구체적으로 6번 문항(AI는 세상의 다양한 문제에 대한 해결책을 제시한다)과 1번 문항(AI는 세상을 더 나은 곳으로 변화시킬 것이다), 그리고 7번 문항(나는 AI를 사용하지 않는 기술을 선호한다)과 4번 문항(AI는 장점보다 단점이 더 많다) 사이의 측정오차 공분산을 허용한

수정 쌍요인 S-1 모형을 구축하여, 모든 지표에서 가장 우수한 적합성을 확인하고 이론적 타당성 또한 높아 수정 쌍요인 S-1 모형을 최종 모형으로 채택하였다. 이러한 결과는 문항 표현 방식에 따른 체계적 반응 편향을 통제할 필요성을 강조하는 기존 연구(최수미, 조영일, 2013; Schriesheim & Hill, 1981; Tepper & Tepper, 1993)와 맥을 같이 한다. 특히 본 연구에서 쌍요인 S-1 모형이 최적의 적합도를 보였다는 점은 긍정적·부정적 표현의 혼합 사용이 내적 구조에 미치는 방법효과를 일반적인 태도(일반요인)와 명확하게 구분하여 평가할 수 있음을 시사한다. 따라서 이 모형을 해석할 때는 척도의 핵심 개념인 일반요인을 중심으로 평가하되, 부정적 표현의 문항들에서는 응답자의 태도와는 무관하게 표현 방식 자체가 응답에 영향을 미치는 별도의 방법효과가 있음을 함께 유념하여 결과를 해석할 필요가 있다.

한편, 인지·정서·행동적 세 가지 측면을 독립된 요인으로 구분한 모형(contents facets)은 원 연구(Stein et al., 2024)와 본 연구 모두에서 적합도가 낮게 나타났다. 따라서 ATTARI-12가 이론적으로는 세 측면을 고려하여 설계되었으나, 실제 응답에서는 이러한 측면들이 명확히 독립되지 않고 일반 태도 요인으로 통합되는 이유에 대한 추가적인 논의가 요구된다. 우선, AI라는 주제 자체가 본질적으로 복잡하고 불확실성이 높은 개념이므로, 응답자들은 각 차원을 구분하기보다 정서적 반응을 보다 직관적으로 표현하는 경향이 존재할 수 있다. 일반적으로 신기술은 응답자의 즉각적인 정서적 반응을 쉽게 촉발하기 때문이다(Park & Woo, 2022; Valor et al., 2022). 특히 AI는 효율성, 기능성뿐만 아니라 따뜻함이나 감정과 같은 관계적 요인이 중요한 사회적 상호작용의 대상

이기 때문에(Merrill et al., 2022), 응답자들이 정서적 평가를 더 강하게 표현했을 수 있다. 또한 태도의 인지·정서·행동적 측면은 이론적으로 독립적이지만 실제 응답자의 심리에서는 명확히 구분하기 어렵다는 점(Rosenberg & Hovland, 1960; Zanna & Rempel, 2008)에서 응답자는 세 가지 요소를 통합된 하나의 태도로 간주하고 응답했을 가능성이 있다. 아울러 ATTARI-12는 설계상 각 차원을 대표하는 문항이 4문항씩 총 12문항으로 제한되어 있다. CFA 문항 구성 원칙상 일반적으로 각 요인당 최소 3개 이상의 문항이 요구되며(Costello & Osborne, 2005; Suhr, 2006), 보다 안정적인 요인을 구성하려면 4개 이상의 문항을 갖추는 것이 권장된다(Costello & Osborne, 2005). 이러한 권장 기준에 비추어 볼 때, 본 척도의 경우 문항 수가 최소한의 권장 기준에는 부합하지만, 보다 견고하고 명확한 차원별 독립성을 확보하기에는 문항 수의 제약으로 인해 한계가 있었을 가능성이 존재한다.

또한, 연구 2에서 실시한 확인적 요인분석 결과, 2번 문항(나는 AI에 대해 부정적인 감정이 강하다), 7번 문항(나는 AI를 사용하지 않는 기술을 선호한다), 8번 문항(나는 AI가 두렵다)의 요인부하량이 .335~.379대 수준으로 낮게 나타났다. 그러나 해당 문항들은 인공지능에 대한 태도의 정서와 행동을 반영하는 내용으로, 본 척도의 이론적 구성 개념상 핵심적인 의미를 지니고 있어 유지하기로 결정하였다. 특히 원 연구(Stein et al., 2024)에서도 이러한 문항들이 포함되어 있으며, 인공지능에 대한 태도의 부정적인 정서 측면을 포괄적으로 측정하기 위해 중요하게 다루어진 바 있다. 또한, 확인적 요인분석에서 요인부하량이 상대적으로 낮게 나타났다고 하더라도, 이론적

타당성과 내용 포괄성을 고려하여 문항을 유지할 수 있다는 선행 연구의 주장을 참고하여 유지하였다(Ximénez, 2009). Ximénez(2009)는 CFA 모형에서 일부 문항의 부하량이 낮게 나타나는 경우에도, 모형이 과도하게 단순화되어 본래 측정하려는 구성개념의 일부가 누락 될 위험이 있기 때문에, 낮은 부하량만을 근거로 문항을 제거하는 것은 바람직하지 않다고 지적한 바 있다. 이에 따라 본 연구에서는 심리측정적 기준뿐만 아니라 이론적 중요성을 고려하여 요인부하량이 낮게 나타난 문항들을 최종 척도에 포함하였다.

셋째, 본 연구에서는 ATTARI-12 척도의 수렴 타당도를 검토하기 위해 기존 인공지능 태도(GAAIS-K), 생성형 인공지능 리터러시(K-GAIL), 이용 의도와의 상관을 분석하였다. 그 결과, ATTARI-12와 GAAIS-K 간에는 높은 정적 상관이 나타났으며, 특히 긍정적 인공지능 태도 요인과의 상관이 가장 높게 확인되었다. 이는 두 척도가 인공지능에 대한 일반적인 태도를 측정한다는 점에서 공통된 구성개념을 반영하고 있음을 보여주며, ATTARI-12의 수렴 타당도를 뒷받침하는 근거로 해석할 수 있다. 다만, 두 척도 간 상관 계수가 .90 이상으로 매우 높은 수준으로 나타났다는 점은, ATTARI-12와 GAAIS-K가 실질적으로 동일한 측정 도구일 가능성에 대한 의문을 불러일으킬 수 있다. 그러나 두 척도는 개발 목적과 대상 측면에서 명확한 차이가 존재한다. ATTARI-12는 직업 여부와 관계없이 성인을 대상으로 일반적인 AI 태도를 측정할 목적으로 개발된 반면, GAAIS는 근로자 집단을 대상으로 직장 내 인공지능 활용 맥락을 염두에 두고 개발된 척도이다(Schepman & Rodway, 2020; Stein et al., 2024). 즉, 측정 대상과 활용

맥락의 차이를 고려할 때, 두 척도는 동일한 도구로 간주하기 어려우며, ATTARI-12는 기존 도구와 유사한 구성개념을 공유하면서도 여러 대상에게 사용할 수 있는 특징을 가진 별개의 측정 도구라고 할 수 있다. 아울러 부정적 태도 요인에서도 ATTARI-12와 GAAIS-K 간 유의한 상관이 확인되어, 본 척도가 인공지능에 대한 태도를 신뢰롭게 측정할 수 있는 도구임을 뒷받침한다.

ATTARI-12는 생성형 인공지능 리터러시 척도(K-GAIL)의 하위 요인 중 AI 활용 능력과 창의적 활용과 유의한 정적 관계를 보였다. 선행 연구에서는 디지털 리터러시가 개인의 디지털 자아효능감을 높이고, 이는 생성형 AI에 대한 긍정적 태도 형성에 영향을 미친다고 보고하였다(김민진 등, 2024; Sumengen et al., 2025). 이러한 선행 연구 결과는 인공지능 활용 능력과 인공지능에 대한 태도 간의 관계가 있다는 본 연구를 지지하는 결과라고 볼 수 있다. 또한 인공지능 이용 의도와도 유의미한 관계를 나타냈는데, 인공지능에 대한 태도가 실제 인공지능 기술을 이용하는 의도에 영향을 미친다고 본 선행 연구 결과와 일치한다(황영훈 등, 2020; Choi et al., 2009). 그러나 K-GAIL의 하위 요인인 비판적 평가와 윤리적 사용 요인과의 관계에서는 유의미한 결과가 나타나지 않았다. 이는 ATTARI-12가 인공지능에 대한 태도 전반은 잘 반영하고 있으나, 인지적 판단과 관련된 구조는 척도 내에 충분히 반영되지 않았을 가능성을 시사한다(김교현, 2002). 특히, 인지, 정서, 행동을 반영한 문항을 구성하였지만, 단일 요인으로 수렴된 점은, 응답자들이 인공지능에 대한 사고적 측면보다는 정서적 반응이나 전반적 평가에 더 강하게 반응했음을 보여주는 결과로 해석될 수 있다.

이러한 결과는 척도의 해석과 활용에 있어 인지적 태도 요소의 한계를 고려할 필요가 있음을 의미하며, 향후 인지적 요인을 보다 명확히 포착할 수 있는 문항 구성에 대한 논의가 요구된다. 이러한 결과를 종합했을 때, ATTARI-12 척도의 수렴 타당도가 확립되었음을 확인할 수 있다.

한편, ATTARI-12의 변별 타당도를 검증하기 위해 컴퓨터 자기효능감(CSEM), 인공지능 불안(AIAS)와의 상관을 분석하였다. ATTARI-12는 CSEM과 낮은 수준의 상관을 보였고, AIAS와는 낮은 부적 상관을 나타냈다. 이는 인공지능에 대한 태도가 기술 활용에 대한 효능감이나 인공지능으로 인한 불안과는 구별된다는 선행 연구 결과와 일치한다(Krägeloh et al., 2024; 박철우, 강경란, 2022). 이러한 결과는 본 척도가 인공지능에 대한 태도를 포괄적으로 측정할 수 있으면서도, 기술에 대한 효능이나 기술로 인해 느끼는 불안과는 명확히 구별됨을 시사한다.

이상의 논의를 바탕으로, 본 연구의 의의는 다음과 같다. 본 연구의 첫 번째 의의는 한국에서 인공지능에 대한 태도를 신뢰롭고 타당하게 측정할 수 있는 도구로서 ATTARI-12 척도의 타당성을 최초로 검증했다는 점이다. 특히 만 19세에서 만 69세에 이르는 폭넓은 연령대와 다양한 직업군을 대상으로 척도를 검증하였기 때문에 척도의 적용 가능성을 높였다는데 의미가 있다.

둘째, ATTARI-12 척도의 타당화는 향후 심리학을 기반으로 한 융합 연구의 발판을 마련하였다는 점에서 의의가 크다. 특히 인공지능 기술의 발전은 글로벌한 현상이며 특정 문화권에 국한되지 않고 보편적인 속성을 지닌다는 점에서, 본 연구에서 ATTARI-12의 요인 구

조, 신뢰도 및 타당도가 원척도와 동일하게 확인된 결과는 이 척도가 한국 문화권에서도 인공지능에 대한 일반적이고 문화 초월적인 태도를 정확하게 반영하고 있음을 시사한다. 이는 향후 인공지능 기술 개발과 정책 수립 과정에서 한국인의 인공지능 수용성을 체계적으로 파악하고, 이를 기반으로 기술의 사회적 적용성을 높이는 데 유용한 근거자료를 제공할 수 있을 것으로 기대된다.

마지막으로, ATTARI-12 척도가 인공지능 활용 능력, 창의적 활용, 이용 의도와 밀접한 관련성을 보인다는 점에서 이 척도의 실용적 활용 가치가 매우 높음을 확인하였다. 따라서 향후 인공지능 관련 교육 및 훈련 프로그램 개발에서 ATTARI-12가 효과적인 평가 도구로 적극 활용될 가능성을 시사한다.

본 연구의 제한점과 후속 연구를 위한 제언은 다음과 같다. 먼저, 본 연구에서는 Stein 등(2024)이 제안한 내용적 하위 요인만을 고려한 모형과 내용적 하위 요인과 부정적 표현의 특수요인을 고려한 모형이 적합도가 충분히 확보되지 않는 한계가 있었다. 이러한 결과는 인지, 정서, 행동 차원의 내용적 하위 요인들이 문항 수준에서 충분히 다뤄지지 않았을 가능성을 시사한다. 따라서 추후 연구에서는 각 하위 요인에 해당하는 문항들이 보다 명확하게 구분되도록 문항 내용을 정교화하고, 필요시 문항 추가 개발과 수정하는 작업이 병행될 필요가 있다. Clark과 Watson(1995)은 내용적 하위 요인이 명확히 드러나지 않는 경우 문항을 새로 작성하거나 문항을 수정하는 과정을 통해 구조적 타당성을 확보할 수 있다고 언급하였다. 즉, 문항을 개발하는 과정에서 인지, 정서, 행동의 하위 요인을 정교하게 측정하는 문항을 추가하거나 수정하여 연구를 확장해

나갈 수 있다.

둘째, 본 연구에서 사용된 실제 표본의 연령 구성이 만 30~40대에 집중되어 있고, 상대적으로 만 60세 이상 연령층이 소수로 나타났다. 따라서 본 연구의 결과를 모든 연령층으로 일반화할 때는 신중한 해석이 필요하다. 향후 연구에서는 연령층별로 표본 수를 균형 있게 확보하여 보다 정확한 일반화 가능성을 검증하는 연구가 필요하다.

셋째, 본 연구는 단회성 온라인 설문조사로 진행된 관계로 검사-재검사 신뢰도를 확보하지 못하였다. 따라서 향후 연구에서는 검사-재검사 신뢰도를 포함한 종단적 접근을 통해 척도의 안정성을 보다 엄밀하게 검토할 필요가 있다. 특히, 대학생 또는 실무자 집단과 같은 특정 집단을 대상으로 한 재검사를 수행하거나, 다양한 표집에서의 반복 검증을 통해 척도의 일반화 가능성을 보다 심도 있게 탐색할 필요가 있겠다.

마지막으로 본 연구는 인공지능에 대한 태도의 일반적 특성을 파악하고자 하였으며, 연구 참여자들이 인공지능을 접하는 경로는 주로 휴대용 기기나 컴퓨터와 같은 일상적 디지털 환경에 제한되어 있었다. 그러나 최근 인공지능 기술은 자율주행차, 지능형 로봇 등 다양한 형태로 빠르게 확산되고 있으며, 이러한 특정 응용 분야에서는 인공지능에 대한 태도가 다르게 나타날 가능성이 있다(Stein et al., 2024). 인공지능 기술이 의료, 교육, 교통 등 여러 영역에 적용되고 있는 현실을 고려할 때, 본 척도가 다양한 맥락에서도 유효하게 작동하는지를 확인하기 위한 추가적인 검증이 필요하다. 따라서 후속 연구에서는 실제적이고 특수한 인공지능 응용 환경을 반영한 타당화 연구를 통해 본 척도의 범용성과 적용 가능성을

을 확장할 것을 제안한다.

참고문헌

- 김교현 (2002). 생명공학에 대한 한국인들의 표상: 대학생들과 일반 성인들을 중심으로. *한국심리학회지: 문화 및 사회문제*, 8(1), 165-187.
- 김남걸 (2001). Likert 척도개발을 위한 문항선정 방법의 비교 분석. [석사학위논문, 연세대학교 대학원].
- 김민진, 김미예, 노환호, 김범수 (2024). 생성형 AI 긍정 태도는 어떻게 형성될 수 있을까? 소비자 경험과 디지털 리터러시가 디지털 자아효능감을 통해 생성형 AI 긍정 태도에 미치는 영향. *소비자학연구*, 35(2), 143-163.
- 김보나, 우종정, 이옥형 (2010). 대학생의 컴퓨터 자기효능감 도구 타당화 및 이러닝 학습효과 간의 관련성 연구. *한국정보기술학회논문지*, 8(9), 153-160.
- 김시내, 박지영 (2024). 인공지능 기술 적용에 대한 한국 직장인의 태도: 척도 개발 및 타당화 연구. *한국심리학회지: 산업 및 조직*, 37(2), 81-117.
- 김아영, 임은영 (2003). 타문화권 척도 변안과정에서 적용되는 절차들 간의 효과 비교. *한국심리학회지: 일반*, 22(1), 89-113.
- 노환호, 김현정, 김민진 (2024). 한국형 생성 인공지능 리터러시 척도 개발 및 타당화. *지식경영연구*, 25(3), 145-171.
- 박선규, 김정순, 정은경 (2024). 대학생이 지각한 인공지능불안이 진로자기효능감에 미치는 영향: 취업 스트레스의 매개효과와

- 계획된 우연기술의 조절효과. *한국심리학회지: 코칭*, 8(2), 59-86.
- 박철우, 강경란 (2022). 대학생들의 자기효능감, 사용경험 및 지각된 유용성이 AI 서비스 수용태도에 미치는 영향 연구. *한국창업학회지*, 17(4), 177-197.
- 서연희, 안정원 (2022). 한국어판 간호대학생의 인공지능에 대한 태도 측정도구 신뢰도와 타당도 검증. *한국간호교육학회지*, 28(4), 357-367.
- 이장주 (2022). 규범의 전환과 사회문제: 코로나를 중심으로. *한국심리학회지: 문화 및 사회문제*, 28(3), 513-527.
- 정민영, 민이수, 임소연 (2023). 돌봄물로서의 인공지능: 20대 청년의 인공지능 불안과 취업스트레스를 중심으로. *과학기술학연구*, 23(2), 112-144.
- 정세정, 신영규 (2024). 디지털전환과 인공지능(AI) 기술에 관한 인식과 태도에 대한 10개국 비교. *보건복지포럼*, 2024(9), 18-36.
- 최수미, 조영일 (2013). 부정문항이 포함된 척도의 요인구조 및 방법효과 검증과 남녀간의 차이 비교: Rosenberg 자기존중감 척도를 중심으로. *한국심리학회지: 일반*, 32(3), 571-589.
- 황영훈, 박수아, 최세정 (2020). 비이용자의 인공지능 스피커에 대한 태도와 행동의도에 영향을 미치는 요인 연구. *미디어 경제와 문화*, 18(1), 31-71.
- 홍세희 (2000). 구조방정식 모형의 적합도 지수 선정기준과 그 근거. *한국심리학회지: 임상*, 19(1), 161-177.
- Ajzen, I., Fishbein, M., Lohmann, S., & Albarracín, D. (2018). The influence of attitudes on behavior. *The handbook of attitudes, volume 1: Basic principles*, 197-255.
- Bochniarz, K., Czerwiński, S., Sawicki, A., & Atroszko, P. (2022). Attitudes to AI among high school students: understanding distrust towards humans will not help us understand distrust towards AI. *Personality and Individual Differences*, 185, 1-7.
- Choi, Y. K., Kim, J., & McMillan, S. J. (2009). Motivators for the intention to use mobile TV: A comparison of South Korean males and females. *International Journal of Advertising*, 28(1), 147-167.
- Clark, L. A., & Watson, D. (1995). Constructing validity: Basic issues in objective scale development. *Psychological Assessment*, 7(3), 309-319.
- Compeau, D. R., & Higgins, C. A. (1995). Computer self-efficacy: Development of a measure and initial test. *MIS Quarterly*, 19(2), 189-211.
- Costello, A. B., & Osborne, J. W. (2005). Best practices in exploratory factor analysis: Four recommendations for getting the most from your analysis. *Practical Assessment, Research and Evaluation*, 10(7), 1-9.
- Curran, P. J., West, S. G., & Finch, J. F. (1996). The robustness of test statistics to nonnormality and specification error in confirmatory factor analysis. *Psychological Methods*, 1(1), 16-29.
- De la Porte, C., Im, Z.J., Sacchi, S., O'Reilly, J., Shin, Y.K., Leschke, J., Citi, M., Ejrnæs, A., Hunt, W., Jensen, M.D., Schütze, C., & Verdin, R. (2024). *Societal Challenges, Public Opinions and Public Policies in 10 countries*

(SCOP-10)

- Fast, E., & Horvitz, E. (2017). Long-term trends in the public perception of artificial intelligence. *Proceedings of the AAAI Conference on Artificial Intelligence*, 31(1).
- Gerlich, M. (2023). Perceptions and acceptance of artificial intelligence: A multi-dimensional study. *Social Sciences*, 12(9), 502.
- 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.
- Hu, L. T., & Bentler, P. M. (1999). Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives. *Structural Equation Modeling: A Multidisciplinary Journal*, 6(1), 1-55.
- 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(7), 1563-1577.
- Kolasinska, A., Lauriola, I., & Quadrio, G. (2019). Do people believe in Artificial Intelligence? A cross-topic multicultural study. In *Proceedings of the 5th EAI International Conference on Smart Objects and Technologies for Social Good* (pp. 31-36).
- Krägeloh, C. U., Melekhov, V., Alyami, M. M., & Medvedev, O. N. (2024). Artificial Intelligence Attitudes Inventory (AIAI): Development and validation using Rasch methodology. *Research Square*.
<https://doi.org/10.21203/rs.3.rs-4403120/v1>
- Merrill, K., Kim, J., & Collins, C. (2022). AI companions for lonely individuals and the role of social presence. *Communication Research Reports*, 39(2), 93-103.
- Nam, T. (2019). Technology usage, expected job sustainability, and perceived job insecurity. *Technological Forecasting and Social Change*, 138, 155-165.
- 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.
- Park, J., Woo, S. E., & Kim, J. (2024). Attitudes towards artificial intelligence application at work: A scale development and validation. *Journal of Occupational and Organizational Psychology*, 97(3), 920-951.
- Quilty, L. C., Oakman, J. M., & Risko, E. (2006). Correlates of the Rosenberg self-esteem scale method effects. *Structural Equation Modeling*, 13(1), 99-117.
- Rosenberg, M. J., & Hovland, C. I. (1960). Cognitive, affective, and behavioral components of attitudes. In C. I. Hovland & M. J. Rosenberg (Eds.), *Attitude organization and change: An analysis of consistency among attitude components*. Yale University Press. (pp. 1-14)
- Schepman, A., & Rodway, P. (2020). Initial validation of the general attitudes towards Artificial Intelligence Scale. *Computers in Human Behavior Reports*, 1, 100014.
- Schriesheim, C. A., & Hill, K. D. (1981). Controlling acquiescence response bias by item reversals: The effect on questionnaire validity.

- Educational and Psychological Measurement*, 41(4), 1101-1114.
- Shank, D. B., Graves, C., Gott, A., Gamez, P., & Rodriguez, S. (2019). Feeling our way to machine minds: People's emotions when perceiving mind in artificial intelligence. *Computers in Human Behavior*, 98, 256-266.
- 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.
- Stein, J., 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.
- Stein, J. P., Messingschlager, T., Gnambs, T., Hutmacher, F., & Appel, M. (2024). Attitudes towards AI: Measurement and associations with personality. *Scientific Reports*, 14(1), 2909.
- Suhr, D. D. (2006). Exploratory or confirmatory factor analysis? In *Proceedings of the 31st Annual SAS Users Group International Conference*. SAS Institute Inc. 200-31, 1-17
- Sumengen, A. A., Subasi, D. O., & Cakir, G. N. (2025). Nursing students' attitudes and literacy toward artificial intelligence: A cross-sectional study. *Teaching and Learning in Nursing*, 20(1), e250-e257.
- Sundar, S. S. (2020). Rise of machine agency: A framework for studying the psychology of human - AI interaction (HAI). *Journal of Computer-Mediated Communication*, 25(1), 74-88.
- Tepper, B. J., & Tepper, K. (1993). The effects of method variance within measures. *The Journal of Psychology: Interdisciplinary and Applied*, 127(3), 293-302.
- Valor, C., Antonetti, P., & Crisafulli, B. (2022). Emotions and consumers' adoption of innovations: An integrative review and research agenda. *Technological Forecasting and Social Change*, 179, 121609.
- Venkatesh, V., & Davis, F. D. (2000). A theoretical extension of the technology acceptance model: Four longitudinal field studies. *Management Science*, 46(2), 186-204.
- 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.
- Ware, J. E. Jr., & Gandek, B. (1998). Methods for testing data quality, scaling assumptions, and reliability: the IQOLA project approach. *Journal of Clinical Epidemiology*, 51(11), 945-952.
- Westerman, D., Edwards, A. P., Edwards, C., Luo, Z., & Spence, P. R. (2020). I-It, I-Thou, I-Robot: The perceived humanness of AI in human-machine communication. *Communication Studies*, 71(3), 393-408.
- Ximénez, C. (2009). Recovery of weak factor loadings in confirmatory factor analysis under conditions of model misspecification. *Behavior Research Methods*, 41(4), 1038-1052.
- Yam, K. C., Tan, T., Jackson, J. C., Shariff, A., & Gray, K. (2023). Cultural differences in

people's reactions and applications of robots, algorithms, and artificial intelligence. *Management and Organization Review*, 19(5), 859-875.

Zanna, M. P., & Rempel, J. K. (2008). Attitudes: A new look at an old concept. In R. H. Fazio & R. E. Petty (Eds.), *Attitudes: Their structure, function, and consequences* (pp. 7-15). Psychology Press.

논문 투고일 : 2025. 04. 01

1 차 심사일 : 2025. 04. 19

게재 확정일 : 2025. 04. 28

Validation of the Korean Version of the Attitudes Toward Artificial Intelligence Scale (ATTARI-12)

Sun Young Lee Hyo Mi Kim Hyun-nie Ahn

Department of Psychology, Ewha Womans University

This study aimed to validate the Korean version of the Attitudes Toward Artificial Intelligence Scale (ATTARI-12), which measures attitudes toward AI across cognitive, emotional, and behavioral dimensions. Although ATTARI-12 was developed to reflect three psychological components, it was originally designed with a unidimensional structure. The scale was translated and administered to 223 Korean adults aged 19 to 69. Exploratory factor analysis showed a two-factor structure separating positively and negatively worded items. To verify this, confirmatory factor analysis was conducted with another independent sample of 223 adults. Competing model analysis was performed, comparing the two-factor model according to the exploratory factor analysis, a bifactor S-1 model reflecting method effects of negative wording, and a modified bifactor S-1 model based on modification indices. The bifactor S-1 model, structurally consistent with the original model, showed the best fit, and the modified version of the bifactor S-1 model was confirmed as the final model. Both convergent and discriminant validity were found to be acceptable. Implications, limitations, and future directions are discussed.

Key words : Korean Version of the Attitudes Toward Artificial Intelligence Scale, Exploratory Factor Analysis, Confirmatory Factor Analysis