

Molnár László – Horváth Kata

Túl a pozitív–negatív dichotómián: egy többkomponensű MI-attitűdskála fejlesztése és validálása

**Beyond the positive-negative dichotomy:
development and validation of a multicomponent AI attitude scale**

Molnár László, a Miskolci Egyetem habilitált egyetemi docense

E-mail: laszlo.molnar@uni-miskolc.hu

Horváth Kata, a Miskolci Egyetem egyetemi tanársegédje

E-mail: kata.horvath@uni-miskolc.hu

A tanulmány célja egy új, az általános MI-attitűdök mérésére alkalmas skála kidolgozása és előzetes validálása. A többkomponensű általános MI-attitűdskála (*Multicomponent Scale of General Attitudes Toward Artificial Intelligence*, MCS–GATAI) a klasszikus háromkomponensű attitűdmodellre épül, de a kognitív, affektív, valamint konatív komponenseken belül külön méri a pozitív és a negatív értékelési irányokat. A skálatejesztés során egy 30 tételes itemkészletet szakértői validálásnak vetettünk alá, majd két kvantitatív adatfelvétel alapján 18 tételes mérőeszközzé redukáltunk. A megerősítő faktorelemzés megfelelő modellilleszkedést jelzett (SRMR = 0,073), a faktorsúlyok minden tétel esetében meghaladták a 0,60-os határértéket, a megbízhatósági mutatók (Cronbach- α : 0,71–0,86; CR: 0,72–0,86) pedig az összes konstrukció esetében jónak tekinthetők. A konvergens validitás ugyancsak megfelelő (AVE: 0,46–0,68), a diszkrimináns validitást pedig a HTMT-mutatók segítségével sikerült igazolnunk. Az ily módon kifejlesztett skála lehetőséget teremt az MI-attitűdök komponensalapú, vegyérték szerint differenciált és potenciálisan ambivalens mintázatainak vizsgálatára.

Kulcsszavak: mesterséges intelligencia, attitűd, skálatejesztés

This study aims to develop and initially validate a new scale for measuring general attitudes towards artificial intelligence. The Multicomponent Scale of General Attitudes Toward Artificial Intelligence (MCS–GATAI) builds on the classical three-component model of attitudes and distinguishes positive and negative evaluative directions within cognitive, affective and conative components. During scale development, an initial pool of 30 items was subjected to expert validation and subsequently reduced to an 18-item instrument based on two quantitative data collections. Confirmatory factor analysis indicated acceptable model fit (SRMR = 0.073), factor loadings exceeded 0.60 for all items, and reliability indicators were adequate across all constructs (Cronbach's α : 0.71–0.86; CR: 0.72–0.86). Convergent validity was generally favourable (AVE: 0.46–0.68), and discriminant validity was clearly supported by HTMT ratios. The scale provides a basis for examining AI attitudes as component-based, valence-differentiated and potentially ambivalent patterns.

Keywords: artificial intelligence, attitude, scale development

A mesterséges intelligencia (MI) az elmúlt évek egyik meghatározó jelenségeként a gazdasági és szervezeti folyamatok mellett a mindennapi élet számos területét is átalakította (Maslej et al., 2025). Az MI-alapú megoldások széles körű terjedése nyomán a technológiához való társadalmi és fogyasztói viszonyulás megértése kiemelt kutatási feladattá vált. A terület kritikusságát erősíti, hogy a mesterséges intelligencia egyszerre értelmezhető a fejlődés, az innováció és a hatékonyságnövelés eszközeként, valamint olyan technológiaként, amely jelentős társadalmi, etikai és pszichológiai dilemmákat vet fel (OECD, 2019). E sajátos kettősség a mesterséges intelligenciával kapcsolatos attitűdök komplexitását növeli, ilyenformán a dichotóm, pozitív–negatív dimenziókra korlátozódó megközelítések egyre kevésbé tekinthetők elégségesnek (Cave et al., 2019; Kelley et al., 2021). Az MI-hez való viszonyulás inkább több, egymással nem feltétlenül harmonizáló attitűdelem együttállásaként, semmint egységes pozitív vagy negatív értékelésként értelmezhető. Az egyének egyszerre ismerhetik el az MI hasznosságát, élhetnek meg vele kapcsolatban bizonytalanságot vagy aggodalmat, és alakíthatnak ki eltérő viselkedési tendenciákat a különböző alkalmazási helyzetekben. Egy felhasználó például elfogadhatja az MI hatékonyságnövelő szerepét, miközben explicit fenntartásai vannak a technológia társadalmi következményeivel kapcsolatban. Hasonlóképpen nyitott lehet bizonyos MI-alapú megoldások kipróbálására, ugyanakkor elutasítóan viszonyulhat a technológia jelenlétéhez az életének érzékenyebb területein. Ez arra utal, hogy az MI-attitűdök vizsgálata olyan differenciált mérési megközelítést igényel, amely az egyidejűleg jelen lévő, akár eltérő irányú viszonyulásokat is képes megragadni.

A probléma különösen méréselméleti szempontból releváns, mivel az általános MI-attitűdök mérésére szolgáló, meglévő skálák csak korlátozottan képesek egyszerre kezelni az attitűdök komponensalapú szerkezetét és a pozitív, illetve a negatív értékelési irányok elkülönítését. Kutatásunk fő mozgatórugója tehát az a felismerés, hogy az általános MI-attitűdök vizsgálatához olyan mérőeszközre van szükség, amely egyszerre képes az attitűdök összetett, háromkomponensű szerkezetének (kognitív – affektív – konatív) kezelésére és az értékelési irányok (pozitív – negatív) elkülönítésére. Meglátásunk szerint nem elegendő az MI-vel kapcsolatos általános vélekedéseket pozitív vagy negatív dimenziók mentén mérni, hanem különbséget kell tenni a pozitív és a negatív kogníciók, az affektív reakciók és a konatív tendenciák között. Egy ilyen 3×2-es mérési logika lehetőséget teremt arra,

hogy az MI-attitűdök ne egyetlen összesített értékelési kontinuumként, hanem strukturált, többdimenziós konstrukcióként jelenjenek meg.

E mérési problémára válaszul dolgoztuk ki a többkomponensű általános MI-attitűdskálát (Multicomponent Scale of General Attitudes Toward Artificial Intelligence, MCS–GATAI). Az MCS–GATAI olyan mérőeszköz, amely a mesterséges intelligenciával kapcsolatos általános attitűdöket komponens és vegyérték szerint differenciált szerkezetben operacionalizálja. A skála hat aldimenziót különít el: (1) pozitív kognitív, (2) negatív kognitív, (3) pozitív affektív, (4) negatív affektív, (5) pozitív konatív és (6) negatív konatív aldimenziót. Ez a felépítés lehetővé teszi, hogy a válaszadók MI-vel kapcsolatos vélekedéseit, érzelmi reakcióit és viselkedési tendenciáit egymástól elkülönülten, ugyanakkor egységes és differenciált attitűdelméleti keretben vizsgálhassuk.

A tanulmány fő célja tehát egy elméletileg megalapozott és empirikusan tesztelt mérőeszköz kidolgozása volt. Az azonosított mérési hiányosságokra és a skálafejlesztési célkitűzésre építve négy egymással összefüggő kutatási célt fogalmaztunk meg:

1. Itemkészlet kidolgozása: az általános MI-attitűdök mérésére alkalmas, elméletileg megalapozott itemkészlet kidolgozása a kognitív, affektív és konatív attitűdkomponensek, valamint a pozitív és a negatív értékelési irányok figyelembevételével.
2. Tartalmi validálás: a kezdeti itemkészlet tartalmi érvényességének vizsgálata szakértői validálás segítségével, különös tekintettel az itemek relevanciájára, érthetőségére és redundanciájára.
3. Strukturális validálás és itemredukció: a skála faktorstruktúrájának és belső szerkezetének empirikus vizsgálata kvantitatív adatfelvételek alapján, valamint a végleges, 18 tételes mérőeszköz kialakítása.
4. Pszichometria validálás: a véglegesített MCS–GATAI-skála előzetes pszichometria validálása a megbízhatóság, a konvergens érvényesség és a diszkrimináns érvényesség mutatói alapján.

E célok együttesen a mesterséges intelligenciával kapcsolatos általános attitűdök mérésének módszertani megalapozását szolgálják. Ennek megfelelően a kutatás egy skálafejlesztési és validálási folyamat keretében valósult meg, amelyben a mérőeszköz kialakítását és empirikus ellenőrzését a skálafejlesztési szakirodalom meghatározó módszertani ajánlásai alapozták meg (*DeVellis, 2017; Schepman–Rodway, 2020; Wang–Wang, 2022*). A kezdeti itemkészlet kialakítását szakértői validálás követte, majd két egymást követő kvantitatív adatfelvétel szolgált a mérőeszköz pszichometria tulajdonságainak vizsgálatára. Az első kvantitatív szakasz célja az itemek előzetes tesztelése és a skála redukciója volt, míg a második kvantitatív szakasz a véglegesített, 18 tételes skálastruktúra validálására irányult. A mérőeszköz értékelése során a faktorstruktúrát, a belső

konzisztenciát, a kompozit megbízhatóságot, valamint a konvergens és a diszkrimináns érvényességet vizsgáltuk.

A tanulmány hozzájárulása három egymáshoz kapcsolódó szinten értelmezhető. Az elméleti hozzájárulása abból áll, hogy az általános MI-attitűdök mérését a klasszikus, háromkomponensű attitűdmodellhez kapcsolja, és rámutat arra, hogy a mesterséges intelligenciához való viszonyulás nem ragadható meg kielégítően egydimenziós vagy pusztán pozitív–negatív értékelési keretben. A kognitív, affektív és konatív komponensek elkülönítése lehetővé teszi az MI-vel kapcsolatos attitűdök belső szerkezetének differenciáltabb értelmezését. Mérésmódszertani szinten a tanulmány egy 3×2-es skálaszerkezetet javasol, amely az egyes attitűdkomponenseken belül külön méri a pozitív és a negatív értékelési irányokat. Empirikus hozzájárulásként a tanulmány bemutatja az MCS–GATAI-skála előzetes validálását magyar felsőoktatási hallgatói mintákon, amely kiindulópontot teremt a mérőeszköz további, nagyobb és heterogénebb mintákon történő teszteléséhez és alkalmazásához. A skála ilyen módon kiindulópontot jelenthet a jövőbeli MI-attitűdkutatások számára, különösen azon vizsgálatok során, amelyek az MI-hez kapcsolódó attitűdök komponensalapú, vegyérték szerint differenciált és potenciálisan ambivalens mintázatait kívánják feltárni.

A tanulmány felépítése a feltárt mérési hiányosságból kiindulva a skáafejlesztés és az előzetes validálás egymásra épülő lépéseit követi. A bevezetést követően egy szakirodalmi áttekintés keretein belül bemutatjuk a háromkomponensű attitűdmodell elméleti alapjait, az általános és kontextusspecifikus MI-attitűdskálák főbb típusait, valamint az attitűdambivalencia méréselméleti jelentőségét (1. fejezet). Ezt követően ismertetjük az MCS–GATAI-skála fejlesztésének és előzetes validálásának folyamatát, beleértve a szakértői validálás, az itemredukció és a két egymást követő kvantitatív adatfelvétel főbb lépéseit (2. fejezet). Az eredményeket bemutató fejezet részletesen tárgyalja a skála pszichometriai értékelését, ezen belül a leíró statisztikákat, a faktorstruktúra vizsgálatát, a megbízhatósági mutatókat, valamint a konvergens és diszkrimináns érvényesség eredményeit (3. fejezet). Végül összegezzük a kutatás főbb elméleti és módszertani tanulságait, bemutatjuk a skála alkalmazhatóságát, valamint kijelöljük a mérőeszköz további validálásának és fejlesztésének lehetséges irányait (4. fejezet).

1. Szakirodalmi áttekintés

A szakirodalmi összefoglaló célja annak bemutatása, hogy az általános MI-attitűdök mérésének problémája miként vezethető le az attitűd klasszikus, háromkomponensű értelmezéséből, a meglévő MI-attitűdskálák, valamint az attitűdam-bivalencia mérési sajátosságaiból. Ennek megfelelően a fejezet elsőként az attitűd fogalmi és komponensalapú értelmezését tárgyalja, majd áttekinti és kritika-ilag értékeli az MI-attitűdök mérésére szolgáló főbb skálátípusokat, végül pedig bemutatja a pozitív és a negatív értékelési irányok elkülönítésének méréselméleti jelentőségét.

1.1. A többkomponensű attitűd értelmezése

A szakirodalomban számos, egymással részben átfedésben álló fogalmi meghatározással találkozhatunk az attitűdkonstrukciót illetően (*Ajzen, 2001; Hamlin, 2016; Katz, 1960*). A különböző megközelítések közös vonásának tekinthető, hogy az attitűdöt valamelyik attitűdobjektummal kapcsolatos egyéni viszonyulásként értelmezik, amelyben az értékítélet központi szerepet tölt be (*Ajzen, 2005; Ajzen–Fishbein, 1980; Edwards, 1957; Osgood et al., 1957*).

A klasszikus attitűdelméleti megközelítések az attitűdöt többkomponensű konstrukcióként értelmezik (*Rosenberg–Hovland, 1960*). E felfogás szerint az attitűd kognitív, affektív és konatív komponensekből épül fel (*Allport, 1954; McGuire, 1985*). A kognitív komponens az attitűd objektummal kapcsolatos vélekedéseket és ismereteket foglalja magában; míg az affektív komponens az érzelmi reakciókra, érzésekre és hangulati viszonyulásokra utal, addig a konatív komponens a cselekvési hajlandóságot, a viselkedési orientációt vagy a használati szándékot ragadja meg (*Haddock–Maio, 2008*). A háromkomponensű attitűdmodell fontos sajátossága, hogy nem feltételezi a komponensek automatikus összhangját, vagyis az attitűdobjektummal kapcsolatos vélekedések, érzelmek és viselkedési tendenciák eltérő irányúak és intenzitásúak lehetnek (*Breckler, 1984; Conner et al., 2021; Zanna–Rempel, 2008*). Ez különösen releváns olyan technológiák esetében, amelyek hatása, hasznossága, általános megítélése alapvetően ellentmondásosnak tekinthető (*Brauner et al., 2025; Cave et al., 2019; Kelley et al., 2021*).

1.2. Az MI-attitűdök mérési megközelítései

A mesterséges intelligenciával kapcsolatos attitűdök mérésére irányuló kutatások jelentős része sok esetben nem tükrözi teljes mértékben ezt az elméleti komplexitást. A meglévő MI-attitűdskálák szisztematikus áttekintése alapján a jelenlegi mérőeszközök alapvetően három nagyobb típusba sorolhatók. Az első csoportba azok a kontextusspecifikus skálák tartoznak, amelyek az MI-vel kapcsolatos attitűdöket valamely konkrét alkalmazási területen, például az oktatásban (Measurement of Attitude in Language Learning with AI, MALL: AI [Yildiz, 2023]; AI-based Scoring Systems Attitude Scale, AI-SAS [Boduroglu–Yigiter, 2026]), az egészségügyben (Artificial Intelligence in Mental Health Scale, AIMHS [Katsiroumpa et al., 2025]), a munkahelyi környezetben (Attitudes Towards Artificial Intelligence at Work, AAAW [Park et al., 2024]), a pszichoterápiában (Artificial Intelligence in Psychotherapy Scale, AIPS [Bilge et al., 2025]) vagy a védelempolitikai alkalmazások (Attitude toward AI in Defense, AAID [Hadlington et al., 2023]) esetében vizsgálják. A kontextusspecifikus mérőeszközök közül jól mutatja a terület differenciálódását a Student Attitude Toward Artificial Intelligence (SATAI-) skála (Suh–Ahn, 2022), amely a tanulók MI-oktatással kapcsolatos attitűdjeinek mérésére készült, és a kognitív, az affektív, valamint a viselkedéses attitűdkomponenseket különíti el. Hasonlóan oktatási fókuszú mérőeszköz az Attitudes Towards Using ChatGPT Scale (ATUC) is (Ajlouni et al., 2023), amely kifejezetten a ChatGPT felsőoktatási tanulási eszközként történő használatával kapcsolatos attitűdöket méri. Ezek a skálák jól illusztrálják, hogy a konkrét alkalmazási környezethez kötött mérőeszközök képesek az adott használati helyzet sajátosságainak érzékenyebb megragadására, ugyanakkor éppen e kontextuális beágyazottság miatt kevésbé alkalmasak az MI-hez mint általános technológiai jelenséghez való viszonyulás vizsgálatára. A második csoportba azok az általános MI-attitűdskálák tartoznak, amelyek elsősorban pozitív és negatív értékelési irányokat különítenek el (például General Attitudes towards Artificial Intelligence Scale, GAAIS [Schepman–Rodway, 2020, 2023]; Short General Attitudes towards Artificial Intelligence Scale, Short GAAIS-10 [Schepman–Rodway, 2026]; Artificial Intelligence Attitudes Inventory, AIAI [Krägeloh et al., 2025]; Attitudes Towards Artificial Intelligence, ATAI [Sindermann et al., 2021]), de nem bontják tovább ezeket kognitív, affektív és konatív komponensekre. Az általános MI-attitűdök mérésére irányuló skálák közül a GAAIS az egyik legfontosabb kiindulópont, mivel általános, nem egyetlen alkalmazási területhez kötött mérőeszközként a pozitív és negatív MI-attitűdök elkülönítésére épül. Hasonló logikát követ a rövidített GAAIS-10, valamint az AIAI is, amelyek szintén pozitív és negatív attitűddimenziók mentén ragadják meg az MI-hez való általános viszonyulást. E mérőeszközök legfőbb előnye, hogy lehetővé teszik az MI-vel kapcsolatos

kedvező és kedvezőtlen értékelések elkülönített mérését, ugyanakkor kevésbé alkalmasak annak feltárására, hogy ezek az értékelések kognitív, affektív vagy konatív szinten jelennek-e meg. A harmadik csoportba azok a mérőeszközök sorolhatók, amelyek elméleti szinten többdimenziós attitűdfelfogásból indulnak ki, de az empirikus validálásuk során a komponensek gyakran aggregált vagy kevésbé differenciált struktúrában jelennek meg (például Attitudes Towards AI, ATTARI-12 [Stein et al., 2024]). Jelen kutatás szempontjából az ATTARI-12 különösen releváns, mivel a skála már egy pszichológiai szempontból megalapozottabb attitűdfelfogásból indul ki: az itemek kialakítása során a szerzők a kognitív, az affektív és a viselkedéses komponenseket, valamint a pozitív és a negatív vegyértékeket is figyelembe vették. Az empirikus validálás eredményei ugyanakkor nem támasztották alá a komponensalapú szerkezet elkülönülését, hanem egy átfogó, egydimenziós, általános MI-attitűdfaktort körvonalaztak, ami jól rámutat a jelen tanulmány által azonosított méréselméleti hiányosságra: továbbra is korlátozottan áll rendelkezésre olyan általános MI-attitűdskála, amely egyidejűleg érvényesíti a háromkomponensű attitűdelméleti logikát és a pozitív, illetve a negatív értékelési irányok komponensen belüli elkülönítését.

1.3. Attitűdambivalencia és méréselméleti következményei

Az MI-attitűdök mérésének egyik központi kérdése, hogy a pozitív és a negatív értékelések egymást kizáró pólusokként vagy egymással párhuzamosan is jelen lévő attitűdelemekként értelmezhetők-e. A mesterséges intelligencia ilyen értelemben kifejezetten alkalmas az ambivalens viszonyulások vizsgálatára, mivel a felhasználók ugyanazon technológiával kapcsolatban egyszerre fogalmazhatnak meg kedvező és kedvezőtlen értékeléseket. Ez a szempont különösen kritikussá teszi a korábban azonosított méréselméleti hiányosságot, mivel az MI-hez való viszonyulás gyakran ambivalens természetű lehet (Dang–Liu, 2021; Fietta et al., 2022; Maier et al., 2020; Mutanga et al., 2024; Pedraja-Rejas et al., 2026; Pinazo-Hernandis–Carcavilla-González, 2026). Az ambivalencia nem egyszerűen bizonytalanságot vagy közömbösséget jelent, hanem pozitív és negatív értékelések egyidejű jelenlétét ugyanazon attitűdobjektummal kapcsolatban (Burgin et al., 2015; Guarana–Hernandez, 2015; Schneider et al., 2021; van Harreveld et al., 2014). Amennyiben egy mérőeszköz a pozitív és a negatív értékeléseket egyetlen kompozit mutatóba vonja össze, az ellentétes irányú válaszok matematikailag kiegyenlíthetik egymást. Ennek következtében az ambivalens attitűdmintázatok láthatatlanná válhatnak. A skálatejesztés szempontjából ezért lényeges, hogy a pozitív és a negatív értékelési irányok ne pusztán ugyanazon kontinuum két végpontjaként, hanem empirikusan elkülöníthető dimenziókként legyenek mérhetők.

2. Módszertan

Az általunk kifejlesztett skála validálása három lépésben történt. Az elsőben szakértői validálást hajtottunk végre, a másodikban és a harmadikban pedig kvantitatív validálásra került sor.

A szakértői validálás 2026 februárjában történt. A kezdeti állításgyűjtemény összesen 30 tételt tartalmazott (mind a hat konstrukcióhoz 5-5 állítást fogalmaztunk meg). Ezt követően felkértünk 7 szakértőt a hazai felsőoktatási intézményekből (3 fő BCE, 1 fő BME, 1 fő BGE, 1 fő PTE, 1 fő SZTE), hogy értékeljék a tételeket egy megadott szempontrendszer alapján. Az értékelési szempontok között az alábbiak szerepeltek: 1. Relevancia (Mennyire reprezentálja az item az adott konstrukciót?), 2. Érthetőség (Érthető-e az item megfogalmazása? Ha nem, hogyan fogalmazná meg?), 3. Redundancia (Átfedésben van-e az item más itemmel? Ha igen, melyikkel?), 4. Javaslat (Megtartaná-e az adott itemet?). A beérkezett értékelések alapján elvégeztük a szükséges korrekciókat, és rátértünk a validálás második szakaszára.

A második szakaszban kvantitatív validálásra került sor, amelyre 2026 márciusának első felében került sor. Online kérdőívet készítettünk, amelyet a Miskolci Egyetem Gazdaságtudományi Kar (ME GTK) hallgatóinak körében osztottunk meg. A beérkezett válaszok tisztítását követően egy 153 fős minta állt a rendelkezésünkre, hogy elvégezzük a megfelelő statisztikai elemzéseket (leíró statisztikák, megerősítő faktorelemzés, megbízhatóság, konvergencia és diszkriminációs validitás). Az elvégzett elemzések alapján az eredetileg 30 tételes skálát 18 tételesre kellett szűkíteniünk (konstrukciónként 3-3 állítás). Az így kapott kérdőívet ezt követően egy újabb tesztelés alá vetettük.

A kérdőív validálásának harmadik szakaszában szintén kvantitatív vizsgálatra került sor, 2026 márciusának második felében. Ezúttal is online kérdőívet használtunk, amelyet az ME GTK első évfolyamos hallgatói körében osztottunk meg. Ennél az adatgyűjtésnél – az adattisztítást követően – egy 93 fős mintához jutottunk, amely a nem véletlen mintavétel miatt nem tekinthető reprezentatívnak.

A kutatás során tesztelt kérdőív az alábbi konstrukciókat és a hozzá tartozó állításokat tartalmazta (1. táblázat).

1. táblázat

A validálás harmadik szakaszában szereplő konstrukciók és állítások
Constructs and items included in the third stage of the validation process

Konstrukciók	Állítások
CP (Cognitive-Positive) Kognitív-Pozitív	CP1 – Az MI képes hatékonyabbá tenni az emberi tevékenységeket.
	CP2 – Az MI segít az összetett információk átláthatóbbá tételében.
	CP3 – A mesterséges intelligencia fejlődése új távlatokat nyit meg az emberiség előtt.
CN (Cognitive-Negative) Kognitív-Negatív	CN1 – Az MI működése túlságosan átláthatatlan és követhetetlen.
	CN2 – Az MI használata hosszú távon leépíti az emberek kritikai gondolkodását.
	CN3 – Az MI-technológiák miatt az emberiség túlságosan kiszolgáltatottá válik a gépeknek.
AP (Affective-Positive) Affektív-Pozitív	AP1 – Lelkesedéssel tölt el az MI-ben rejlő lehetőségek látványa.
	AP2 – Kifejezetten élvezem az MI-vel való interakció modern élményét.
	AP3 – Optimista vagyok az MI mindennapi életre gyakorolt hatásával kapcsolatban.
AN (Affective-Negative) Affektív-Negatív	AN1 – Félek attól, hogy az MI fejlődése kicsúszik az emberi kontroll alól.
	AN2 – Bizalmatlan vagyok minden olyan szolgáltatással szemben, amely MI-t használ.
	AN3 – Nyomasztónak találok az MI-technológiák gyors és megállíthatatlan terjedését.
BP (Behavioural-Positive) Konatív-Pozitív	BP1 – Szívesen integrálok MI-alapú megoldásokat a napi rutinom közé.
	BP2 – Másoknak is javasolnám, hogy használjanak MI-t a feladataik megkönnyítésére.
	BP3 – Nyitott vagyok arra, hogy MI-vel működjek együtt egy közös feladat megoldásában.
BN (Behavioural-Negative) Konatív-Negatív	BN1 – Tudatosan kerülöm azokat a helyzeteket, ahol MI-t kellene használnom.
	BN2 – Ellenállok annak, hogy az MI-t bevezessék az életem fontos területeire.
	BN3 – Igyekszem korlátozni az MI jelenlétét a közvetlen környezetemben.

Forrás: saját szerkesztés.

A fenti állításokat hétfokú Likert-skála segítségével kellett a válaszadóknak értékelniük, ahol az 1-es „egyáltalán nem ért egyet”, a 7-es pedig „kifejezetten egyetért” értelemben szerepelt. A beérkezett válaszokat ezt követően adattisztításnak vetettük alá, amely során töröltük a hiányzó értékeket tartalmazó rekordokat, valamint azokat, amelyek alacsony szórást (pl. csak 4-es értéket adtak) mutattak. Így jutottunk el ahhoz a 93 fős mintához, amelynek az SPSS-ben és az ADANCO-ban készült elemzési eredményeit a következő fejezetben ismertetjük.

3. Eredmények

3.1. Az elemzések áttekintése

A skála kvantitatív validálása során több lépcsős statisztikai elemzést végeztünk a mérőeszköz pszichometriai jellemzőinek feltárása érdekében. Az elemzési folyamat az adatok alapvető jellemzőinek feltérképezésével indult, beleértve a leíró statisztikákat és a multikollinearitás vizsgálatát, hogy igazoljuk az adatsor alkalmazását a további többváltozós próbákhoz. A skála strukturális érvényességét megerősítő faktorelemzéssel (Confirmatory Factor Analysis, CFA) teszteltük, ahol az illeszkedés jóságát többek között a Standardized Root Mean Residual (SRMR-) mutatóval ellenőriztük. A struktúra rögzítését követően a mérőeszköz megbízhatóságát a tradicionális Cronbach-alfa, valamint a modernebb kompozit megbízhatósági mutató (Composite Reliability, CR) segítségével értékeltük. Végezetül a skála validitásának mélyebb elemzése céljából megvizsgáltuk a konvergens validitást az átlagos kivonatolt variancia (Average Variance Extracted, AVE) mentén, valamint a diszkrimináns validitást, amelyet a klasszikus Fornell–Larcker-kritérium és a szigorúbb Heterotrait-Monotrait (HTMT) eljárás alkalmazásával támasztottunk alá. Az alábbiakban ezen elemzések részletes eredményeit mutatjuk be.

3.2. Leíró statisztikák és az adatok eloszlása

Az eredmények bemutatását a skála tételeinek (itemek) leíró elemzésével kezdtük. A 2. táblázat tartalmazza az egyes itemek átlagait, varianciáit, valamint a ferdeségi (*skewness*) és csúcsossági (*kurtosis*) mutatókat.

2. táblázat

A mérési tételek leíró statisztikái
Descriptive statistics of the measurement items

Konstrukció	Állítás kódja	Átlag	Variancia	Ferdeség	Csúcsosság
Kognitív-Pozitív	CP1	5,52	1,10	–0,42	–0,46
	CP2	5,59	1,51	–0,65	–0,31
	CP3	5,48	1,54	–0,79	1,05
Kognitív-Negatív	CN1	3,24	2,01	0,41	–0,61
	CN2	5,51	1,93	–0,67	–0,44
	CN3	4,90	2,11	–0,28	–0,92
Affektív-Pozitív	AP1	4,35	2,19	–0,22	–0,29
	AP2	4,39	2,20	–0,02	–0,55
	AP3	4,15	1,93	0,10	–0,46
Affektív-Negatív	AN1	4,30	3,15	–0,06	–1,15
	AN2	3,58	2,55	0,46	–0,59
	AN3	4,03	3,29	0,11	–1,19
Konatív-Pozitív	BP1	4,19	2,51	–0,14	–0,58
	BP2	4,84	2,01	–0,15	–0,47
	BP3	5,16	1,92	–0,80	0,41
Konatív-Negatív	BN1	2,71	2,34	1,06	0,73
	BN2	3,41	2,72	0,46	–0,53
	BN3	3,60	2,33	0,14	–0,69

Forrás: saját szerkesztés.

Az átlagértékek és a varianciák vizsgálata alapján az itemek megfelelően differenciálnak a válaszadók között, nem tapasztaltunk szélsőséges „plafon-” vagy „padlóeffektust”. A normalitás vizsgálata kulcsfontosságú volt a konfirmatív faktorelemzés becslési eljárásának kiválasztásához. A ferdeség értékei egy kivétellel minden esetben a $-1,0$ és $+1,0$ tartományon belül maradtak, míg a csúcsossági mutatók nem haladták meg a $\pm 2,0$ kritikus értéket. Mivel a mutatók nem jeleztek jelentős eltérést a normál eloszlástól, az adatsor alkalmasnak bizonyult a CFA elvégzésére maximum likelihood (ML-) módszerrel.

A mérőeszköz szerkezeti elemzése előtt megvizsgáltuk a változók közötti multikollinearitás esetleges jelenlétét. Ennek ellenőrzésére az ún. *Variance inflation factor* (VIF) mutatót alkalmaztuk. A multikollinearitás fennállása torzíthatja a faktorsúlyok becslését, azonban az eredmények alapján ez a veszély nem áll fenn.

3. táblázat

A mérési tételek multikollinearitás-vizsgálata
Multicollinearity assessment of the measurement items

Konstrukció	Állítás kódja	VIF
Kognitív-Pozitív	CP1	2,30
	CP2	1,67
	CP3	2,31
Kognitív-Negatív	CN1	1,29
	CN2	1,45
	CN3	1,63
Affektív-Pozitív	AP1	2,60
	AP2	2,27
	AP3	1,53
Affektív-Negatív	AN1	1,92
	AN2	1,38
	AN3	2,19
Konatív-Pozitív	BP1	2,07
	BP2	2,36
	BP3	2,26
Konatív-Negatív	BN1	2,05
	BN2	2,08
	BN3	2,35

Forrás: saját szerkesztés.

A 3. táblázatban látható, hogy a VIF-értékek az 1,29 és 2,60 tartományban mozognak. Mivel minden érték jelentősen elmarad a szakirodalomban kritikusnak tekintett 5-ös (vagy szigorúbb megközelítés esetén 3,3-es) küszöbértéktől, megállapítható, hogy az itemek közötti redundancia nem ér el olyan mértéket, amely korlátozná a további többváltozós elemzések megbízhatóságát.

3.3. Konstruktumvaliditás: megerősítő faktorelemzés

A mérőeszköz strukturális érvényességének igazolására konfirmatív faktorelemzést végeztünk. A modell illeszkedésének vizsgálatakor az SRMR-mutatót vettük figyelembe, amelynek értéke 0,0730 lett. Mivel ez az érték elmarad a szak-

irodalomban elfogadott 0,08-os küszöbértéktől (*Hu–Bentler, 1999*), megállapítható, hogy a feltételezett modellstruktúra megfelelően illeszkedik a tapasztalati adatokhoz.

A modell illeszkedésén túlmenően megvizsgáltuk az egyes itemek faktoron belüli súlyait is. A 4. táblázat mutatja be a sztenderdizált faktorsúlyokat (*factor loadings*).

4. táblázat

Megbízhatósági és érvényességi statisztikák
Reliability and validity statistics

Konstrukció	Állítás kódja	Faktorsúly	Cronbach-féle alfa	CR	AVE
Kognitív-Pozitív	CP1	0,81	0,84	0,84	0,63
	CP2	0,84			
	CP3	0,73			
Kognitív-Negatív	CN1	0,61	0,71	0,72	0,46
	CN2	0,73			
	CN3	0,69			
Affektív-Pozitív	AP1	0,75	0,82	0,82	0,61
	AP2	0,81			
	AP3	0,77			
Affektív-Negatív	AN1	0,72	0,78	0,79	0,55
	AN2	0,68			
	AN3	0,82			
Konatív-Pozitív	BP1	0,84	0,86	0,86	0,68
	BP2	0,77			
	BP3	0,86			
Konatív-Negatív	BN1	0,79	0,86	0,86	0,67
	BN2	0,81			
	BN3	0,85			

Forrás: saját szerkesztés.

Ahogy a 4. táblázatban látható, minden tétel faktorsúlya meghaladja a kritikusnak tekintett 0,60 értéket, a legtöbb item esetében pedig 0,70 feletti töltődést tapasztaltunk. Az adatok alapján a tételek statisztikailag szignifikáns módon és megfelelő erősséggel kapcsolódnak a hozzájuk rendelt látens változókhoz, ami alátámasztja a skála konstruktumvaliditását.

3.4. Megbízhatóság

A mérőeszköz belső konzisztenciáját két megközelítésben vizsgáltuk. Egyrészt kiszámítottuk a tradicionális Cronbach-alfa mutatókat, másrészt a modernebb és pontosabb becslést adó kompozit megbízhatóságot (Jöreskog-féle mutató). (Az eredményeket szintén a 4. táblázatban láthatjuk.) A Cronbach-alfa értékek a 0,71 és 0,86 közötti tartományban mozognak, ami a szakirodalmi ajánlások (Nunnally, 1994) szerint jó belső konzisztenciát jelez, mivel minden faktor esetében meghaladják a 0,70-os küszöbértéket. Mivel a Cronbach-alfa hajlamos alulbecsülni a megbízhatóságot a CFA-modellek esetében, kiemelt figyelmet fordítottunk a kompozit megbízhatóság mutatóira. A CR-értékek 0,72 és 0,86 között alakultak, ami tovább erősíti a skála stabilitását. Tekintettel arra, hogy mindkét mutató konzisztensen a 0,70-os határérték felett van, megállapítható, hogy a mérőeszköz faktorai megbízhatóan és stabilan mérik a vizsgált látens konstrukciókat.

3.5. Konvergens és diszkrimináns érvényesség

A validálási folyamat zárásaként a faktorok közötti összefüggéseket és azok egyedi jellegét vizsgáltuk. A konvergens validitást az átlagos kivonatolt variancia mutatóival ellenőriztük. Az AVE-értékek a 0,46 és 0,68 közötti tartományban mozognak (4. táblázat). Bár az egyik faktor (CN) esetében az érték némileg elmarad az ideális 0,50-os küszöbtől, mivel a hozzá tartozó kompozit megbízhatóság (CR) minden esetben meghaladja a 0,70-ot, a konstrukció konvergens validitása elfogadhatónak tekinthető (Fornell–Larcker, 1981).

A diszkrimináns validitás vizsgálatához két módszert alkalmaztunk. Elsőként a Fornell–Larcker-kritériumot ellenőriztük, amely szerint az AVE négyzetgyökének nagyobbnak kell lennie a többi látens változóval való korrelációnál. Az eredmények (5. táblázat) alapján a feltétel egy kivétellel teljesült.

5. táblázat

Diszkrimináns érvényesség a Fornell–Larcker-kritérium alapján
Discriminant validity based on the Fornell–Larcker criterion

	CP	CN	AP	AN	BP	BN
CP	0,63					
CN	0,06	0,46				
AP	0,19	0,04	0,61			
AN	0,06	0,61	0,11	0,55		
BP	0,33	0,04	0,47	0,14	0,68	
BN	0,18	0,26	0,29	0,38	0,48	0,67

Forrás: saját szerkesztés.

6. táblázat

Diszkrimináns érvényesség a HTMT-arány alapján
Discriminant validity based on the HTMT Ratio

	CP	CN	AP	AN	BP	BN
CP						
CN	0,23					
AP	0,44	0,21				
AN	0,25	0,78	0,34			
BP	0,57	0,21	0,68	0,37		
BN	0,42	0,52	0,54	0,62	0,69	

Forrás: saját szerkesztés.

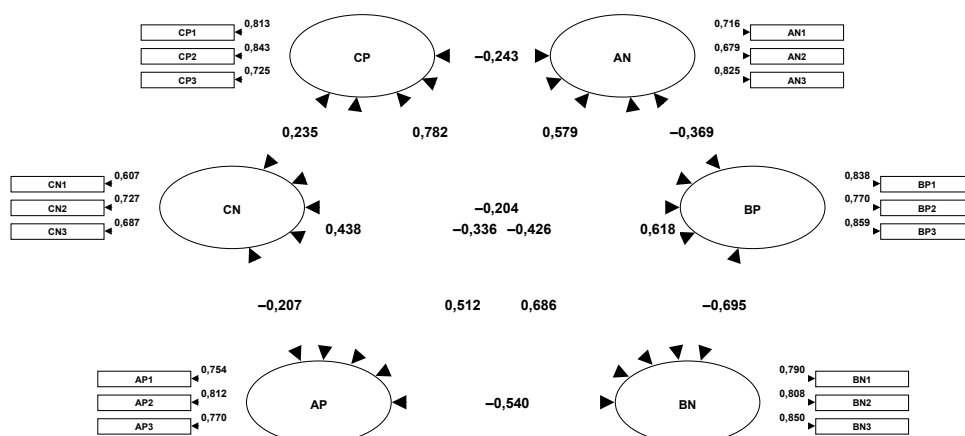
Végezetül elvégeztük a szigorúbb HTMT-elemzést is. Az összes kapott HTMT-értékek a kritikus 0,85-os határ alatt maradtak, ami egyértelműen igazolja, hogy a skála dimenziói egymástól jól elkülöníthető, önálló fogalmakat mérnek.

3.6. Az elemzések összefoglalása

Az eredmények vizuális áttekintése érdekében az 1. ábra szemlélteti a véglegesített mérési modellt. Az ábrán látható útvonalak és korrelációs értékek megerősítik a konstrukciók közötti elméletileg megalapozott kapcsolatokat.

1. ábra

A validált mérési modell struktúrája és a konstrukciók közötti kapcsolatok
Structure of the validated measurement model and relationships among the constructs



Forrás: saját szerkesztés.

Az ábrán látható sztenderdizált értékek összhangban vannak a korábban bemutatott diszkriminációs validitási mutatókkal, és egyértelműen szemléltetik a skála belső integritását. Ezzel a mérőeszköz kvantitatív validálása lezárult, igazolva a skála alkalmasságát a további felhasználásra.

4. Következtetések

A tanulmány célja egy olyan új mérőeszköz kidolgozása és előzetes validálása volt, amely a mesterséges intelligenciával kapcsolatos általános attitűdöket a szakirodalomban általánosan megjelenő egydimenziós megközelítés helyett komponens és vegyérték szerint differenciált szerkezetben ragadja meg. A kiindulópontot az a méréselméleti probléma jelentette, hogy a meglévő MI-attitűdskálák jelentős része vagy konkrét alkalmazási kontextushoz kötődik, vagy az általános MI-attitűdöket elsősorban pozitív és negatív értékelési irányok mentén méri. Ezen megközelítések nem teszik lehetővé a kognitív, affektív és konatív komponensek elkülönített vizsgálatát, az ellentétes vélekedések differenciált megértését és az ambivalencia adekvát feltárását. E hiányosságra válaszul dolgoztuk ki a többkomponensű általános MI-attitűdskálát (*Multicomponent Scale of General Attitudes Toward Artificial Intelligence*, MCS–GATAI), amely 6 aldimenzió mentén operacionálizálja az MI-hez való felhasználói viszonyulást.

Eredményeink összességében alátámasztják a javasolt 3×2-es mérési szerkezet alkalmazhatóságát. A szakértői validálás hozzájárult az itemkészlet tartalmi pontosításához, míg a két egymást követő kvantitatív adatfelvétel lehetővé tette a skála empirikus finomítását és pszichometriai értékelését. Az eredetileg 30 tételes itemkészlet a validálási folyamat során 18 tételes mérőeszközzé redukálódott, amely konstrukciónként 3-3 állítást tartalmaz. A véglegesített skálaszerkezet megerősítő faktorelemzése elfogadható modellilleszkedést mutatott, az itemek faktorsúlyai megfelelőek voltak, a belső konzisztenciát jelző Cronbach-alfa és a kompozit megbízhatósági mutatók pedig minden konstrukció esetében elérték az elfogadott küszöbértékeket. A konvergens validitás eredményei alapvetően kedvezőek, bár a negatív kognitív dimenzió esetében az AVE értéke némileg elmaradt az ideális 0,50-os határértéktől. Ez a skála későbbi továbbfejlesztése során külön figyelmet igényel, ugyanakkor a megfelelő kompozit megbízhatósági érték alapján a konstrukció előzetes megtartása indokoltnak tekinthető. A diszkriminációs érvényesség vizsgálata, különösen a HTMT-mutatók alapján, azt támasztotta alá, hogy a skála dimenziói empirikusan is elkülöníthető fogalmi egységeket alkotnak.

A kutatás elméleti hozzájárulása az, hogy az általános MI-attitűdök mérését a klasszikus háromkomponensű attitűdelméleti megközelítéshez kapcsolja. A kognitív, affektív és konatív komponensek elkülönítése különösen indokolt a mester-séges intelligencia esetében, mivel az MI-hez való viszonyulás gyakran nem egy-séges, hanem belsőleg differenciált és potenciálisan ambivalens szerkezetű. Ez összhangban áll azokkal az attitűdelméleti megközelítésekkel, amelyek szerint az attitűdobjektummal kapcsolatos vélekedések, érzelmi reakciók és viselkedési tendenciák nem szükségszerűen rendeződnek azonos módon (*Breckler, 1984; Conner et al., 2021; Zanna–Rempel, 2008*). Jelen kutatás eredményei ezt a logi-kát az MI-attitűdök mérésének területére ültetik át, és egyúttal bizonyítják, hogy az MI-hez való általános viszonyulás differenciáltabb módon ragadható meg, ha az attitűdkomponensek és az értékelési irányok egyszerre jelennek meg a mérési struktúrában.

A tanulmány mérés módszertani hozzájárulása az, hogy a meglévő általános MI-attitűdskálákhoz képest differenciáltabb, komplex mérőeszközt javasol. A GAAIS és annak rövidített változata fontos előrelépést jelentettek azáltal, hogy a pozitív és a negatív MI-attitűdöket elkülönítetten kezelték (*Schepman–Rodway, 2020, 2023, 2026*). Hasonló logika jelenik meg az AIAI esetében is, amely szintén a pozitív és a negatív attitűddimenziók elkülönítésére épít (*Krägeloh et al., 2025*). Ezek a mérőeszközök ugyanakkor kevésbé teszik lehetővé annak feltárását, hogy a pozitív és a negatív értékelések milyen attitűdkomponensekhez kapcsolódnak. Az ATTARI-12 ezzel szemben már pszichológiailag megalapozottabb attitűdfel-fogásból indul ki, mivel az itemek kialakítása során figyelembe veszi a kognitív, az affektív és a viselkedéses elemeket, valamint a pozitív és a negatív vegyértéket is (*Stein et al., 2024*). Ugyanakkor a szerzők ezt az összetett struktúrát empirikusan nem tudták igazolni. Az MCS–GATAI ehhez képest abban kínál új megközelítést, hogy az általános MI-attitűdök mérésében nem olvasztja össze a különböző kom-ponenseket egyetlen dimenzióba, hanem a pozitív és a negatív kogníciókat, affek-tív reakciókat és konatív tendenciákat külön-külön méri.

A skála egyik fontos előnye, hogy alkalmas lehet az MI-vel kapcsolatos atti-tűdambivalencia differenciáltabb vizsgálatára. Az ambivalencia szakirodalma sze-rint a pozitív és a negatív értékelések egyidejű jelenléte nem tekinthető pusztán közömbösségnek vagy bizonytalanságnak, hanem önálló attitűdszerkezeti jelen-séggként értelmezhető (*Ajzen, 2001; Burgin et al., 2015; Schneider et al., 2021; van Harreveld et al., 2014*). Az MCS–GATAI-skála ebben a tekintetben azért lehet hasznos, mert nem csupán azt teszi mérhetővé, hogy a válaszadók pozitívan vagy negatívan viszonyulnak-e az MI-hez, hanem azt is, hogy ez a pozitív vagy negatív viszonyulás elsősorban kognitív szinten, a felhasználó érzelmi reakcióiban vagy cselekvéseiben jelenik-e meg. Ez különösen fontos lehet olyan esetekben, amikor a válaszadó egyszerre ismeri el az MI hasznosságát, ugyanakkor érzelmi szinten

aggodalmat él meg, vagy amikor kedvező kognitív értékelései ellenére elutasító viselkedési orientációt mutat bizonyos MI-alapú alkalmazásokkal szemben.

Mindezekén túl a skála kiindulópontot teremthet a mesterséges intelligenciával kapcsolatos általános attitűdök részletesebb vizsgálatához is. A mérőeszköz alkalmas lehet különböző társadalmi csoportok, fogyasztói szegmensek, hallgatói vagy munkavállalói csoportok MI-attitűdjeinek összehasonlítására, valamint annak feltárására, hogy az MI-hez való viszonyulás milyen komponensek mentén mutat pozitív vagy negatív mintázatot. A skála különösen hasznos lehet olyan kutatásokban, amelyek nem egyetlen konkrét MI-alkalmazás elfogadását kívánják mérni, hanem az MI-hez mint szélesebb technológiai és társadalmi jelenséghez való általános viszonyulás differenciált megértésére törekcszenek. Ilyen módon az MCS–GATAI kiegészítheti a kontextusspecifikus mérőeszközöket, például az oktatási, munkahelyi, egészségügyi vagy pszichoterápiás környezetben alkalmazott MI-attitűdskálákat, miközben lehetőséget ad az általánosabb, kontextusokon átívelő attitűdmintázatok feltárására is.

Ugyanakkor a kedvező eredmények mellett a kutatás több korláttal is rendelkezik. Elsőként fontos hangsúlyoznunk, hogy a tanulmányban ismertetett eredmények a skála kezdeti validálását szolgálták, az elvégzett kvantitatív vizsgálatok nem reprezentatív, felsőoktatási hallgatói mintákon történtek. Ennek megfelelően az eredmények nem általánosíthatók. Másrészt, bár a pszichometriai mutatók összességében kedvezőnek tekinthetők, a negatív kognitív dimenzió konvergens validitása további vizsgálatot igényel. Továbbá a kutatás keresztmetszeti jellege miatt nem alkalmas annak feltárására, hogy az MI-hez kapcsolódó attitűdkomponensek miként változnak a használati tapasztalatok, a technológia fejlődése vagy a társadalmi diskurzus átalakulása következtében. Végezetül korlátként azonosítható, hogy a skála jelenlegi formája önbevalláson alapul, így az eredményeket befolyásolhatják a válaszadói torzítások, például a társadalmi elvárásokhoz való igazodás vagy az MI-vel kapcsolatos aktuális közbeszéd hatása is.

Jelen tanulmány alapján több olyan kutatási irányt is azonosíthatunk, amelyek a mérőeszköz további validálása, finomítása és szélesebb körű alkalmazása szempontjából meghatározók lehetnek. Elsőként indokoltnak látjuk a skála nagyobb, heterogénebb és reprezentatívabb mintákon történő további validálását, méghozzá különböző életkori, iskolázottsági, foglalkozási és kulturális csoportok bevonásával. Másrészt célszerűnek tartjuk annak vizsgálatát, hogy a skála mérési invarianciája miként alakul, különösen nem, életkor, MI-használati tapasztalat, képzési terület vagy munkakör szerinti bontásban. Továbbá a skála prediktív érvényességének vizsgálata is fontos következő lépés lehet: ennek keretében feltárható, hogy az egyes attitűdkomponensek milyen mértékben kapcsolódnak az MI használati gyakoriságához, elfogadásához, elutasításához, a bizalom kialakulásához vagy a koc-

kázatészleléshez. Végezetül, érdemes lenne longitudinális kutatásokban is alkalmazni a mérőeszközt annak vizsgálatára, hogy az MI-vel kapcsolatos kognitív, affektív és konatív attitűdelemek hogyan alakulnak a technológiával szerzett tapasztalatok bővülése által.

Összességében a tanulmány egy olyan új, elméletileg megalapozott és empirikusan validált mérőeszközt mutat be, amely hozzájárulhat a mesterséges intelligenciával kapcsolatos általános attitűdök differenciáltabb méréséhez és értelmezéséhez. Az MCS–GATAI-skála legfontosabb újdonsága, hogy egyszerre érvényesíti a háromkomponensű attitűdelméleti logikát és a pozitív, illetve a negatív értékelési irányok elkülönítését. Ezáltal a mérőeszköz lehetőséget teremt arra, hogy az MI-attitűdöket ne leegyszerűsített pozitív–negatív kontinuumként, hanem többdimenziós, komponensalapú és potenciálisan ambivalens attitűdszerkezetként vizsgálhassuk. A skála további validálása és alkalmazása hozzájárulhat ahhoz, hogy a mesterséges intelligenciával kapcsolatos társadalmi, fogyasztói és felhasználói viszonyulásokat a jövőben árnyaltabb, elméletileg megalapozottabb és módszertanilag pontosabb módon mérjük.

*

Jelen tanulmány a Kulturális és Innovációs Minisztérium Egyetemi Kutatói Ösztöndíj Programjának a Nemzeti Kutatási, Fejlesztési és Innovációs Alapból finanszírozott szakmai támogatásával készült.

Irodalom

- Ajlouni, A. – Wahba, F. – Almahaireh, A. (2023): Students' Attitudes Towards Using ChatGPT as a Learning Tool: The Case of the University of Jordan. *International Journal of Interactive Mobile Technologies (IJIM)*, 17(18), 99–117. <https://doi.org/10.3991/ijim.v17i18.41753>
- 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>
- Ajzen, I. (2005): *Attitudes, personality, and behavior (2nd ed)*. UK Higher Education OUP Psychology Psychology. McGraw-Hill Education. <https://books.google.hu/books?id=dmJ9EGEy0ZYC>
- Ajzen, I. – Fishbein, M. (1980): *Understanding Attitudes and Predicting Social Behavior (1st ed)*. Prentice-Hall. <https://books.google.hu/books?id=AnNqAAAAMAAJ>
- Allport, G. W. (1954): *The Nature of Prejudice*. Addison-Wesley Publishing Company. <https://books.google.hu/books?id=u94XUyRuDI4C>
- Bilge, Y. – Emiral, E. – Delal, İ. (2025): Attitudes towards Artificial intelligence (AI) in psychotherapy: Artificial Intelligence in Psychotherapy Scale (AIPS). *Clinical Psychologist*, 29(2), 194–204. <https://doi.org/10.1080/13284207.2025.2515074>
- Boduroglu, E. – Yigiter, M. S. (2026): Artificial intelligence scoring attitudes: Scale development and validation. *Education and Information Technologies*, 31(3), 701–726. <https://doi.org/10.1007/s10639-025-13836-7>

- Brauner, P. – Glawe, F. – Liehner, G. L. – Vervier, L. – Ziefle, M. (2025): Mapping public perception of artificial intelligence: Expectations, risk–benefit tradeoffs, and value as determinants for societal acceptance. *Technological Forecasting and Social Change*, 220, 124304. <https://doi.org/10.1016/j.techfore.2025.124304>
- Breckler, S. J. (1984): Empirical validation of affect, behavior, and cognition as distinct components of attitude. *Journal of Personality and Social Psychology*, 47(6), 1191–1205. <https://doi.org/10.1037/0022-3514.47.6.1191>
- Burgin, C. J. – Chun, C. A. – Horton, L. E. – Barrantes-Vidal, N. – Kwapil, T. R. (2015): Splitting of Associative Threads: The Expression of Schizotypal Ambivalence in Daily Life. *Journal of Psychopathology and Behavioral Assessment*, 37(2). <https://doi.org/10.1007/s10862-014-9457-7>
- Cave, S. – Coughlan, K. – Dihal, K. (2019): „Scary Robots”: Examining Public Responses to AI. *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*, 331–337. <https://doi.org/10.1145/3306618.3314232>
- Conner, M. – Wilding, S. – Van Harreveld, F. – Dalege, J. (2021): Cognitive-Affective Inconsistency and Ambivalence: Impact on the Overall Attitude–Behavior Relationship. *Personality and Social Psychology Bulletin*, 47(4), 673–687. <https://doi.org/10.1177/0146167220945900>
- Dang, J. – Liu, L. (2021): Robots are friends as well as foes: Ambivalent attitudes toward mindful and mindless AI robots in the United States and China. *Computers in Human Behavior*, 115, 106612. <https://doi.org/10.1016/j.chb.2020.106612>
- DeVellis, R. F. (2017): *Scale development: Theory and applications (Fourth edition)*. SAGE. https://www.mpm.ac.in/News/Ebook_mpm_sadmin07755388.pdf
- Edwards, A. L. (1957): *Techniques of Attitude Scale Construction*. Appleton-Century-Crofts. <https://doi.org/10.1037/14423-000>
- Fietta, V. – Zecchinato, F. – Stasi, B. D. – Polato, M. – Monaro, M. (2022): Dissociation Between Users’ Explicit and Implicit Attitudes Toward Artificial Intelligence: An Experimental Study. *IEEE Transactions on Human-Machine Systems*, 52(3), 481–489. <https://doi.org/10.1109/THMS.2021.3125280>
- Fornell, C. – Larcker, D. F. (1981): Evaluating Structural Equation Models with Unobservable Variables and Measurement Error. *Journal of Marketing Research*, 18(1). <https://doi.org/10.2307/3151312>
- Guarana, C. L. – Hernandez, M. (2015): Building sense out of situational complexity: The role of ambivalence in creating functional leadership processes. *Organizational Psychology Review*, 5(1). <https://doi.org/10.1177/2041386614543345>
- Haddock, G. – Maio, G. R. (2008): *Attitudes: Content, Structure and Function*. In: M. Hewstone – W. Stroebe – K. Jonas: *Introduction to social psychology: A European perspective* (4th ed, 112. p.). BPS Blackwell.
- Hadlington, L. – Binder, J. – Gardner, S. – Karanika-Murray, M. – Knight, S. (2023): The use of artificial intelligence in a military context: Development of the attitudes toward AI in defense (AAID) scale. *Frontiers in Psychology*, 14, 1164810. <https://doi.org/10.3389/fpsyg.2023.1164810>
- Hamlin, R. (2016): Functional or constructive attitudes: Which type drives consumers’ evaluation of meat products? *Meat Science*, 117, 97–107. <https://doi.org/10.1016/j.meatsci.2016.02.038>
- Hu, L. – 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. <https://doi.org/10.1080/10705519909540118>

- Katsiroumpa, A. – Konstantakopoulou, O. – Moisoglou, I. – Gallos, P. – Galani, O. – Lialiou, P. – Tsiachri, M. – Galanis, P. (2025): Development and Validation of the Artificial Intelligence in Mental Health Scale: *Application for AI Mental Health Chatbots. Healthcare*, 13(24), 3269. <https://doi.org/10.3390/healthcare13243269>
- Katz, D. (1960): The Functional Approach to the Study of Attitudes. *The Public Opinion Quarterly*, 24(2), 163–204. JSTOR., <https://doi.org/10.1086/266945>
- Kelley, P. G. – Yang, Y. – Heldreth, C. – Moessner, C. – Sedley, A. – Kramm, A. – Newman, D. T. – Woodruff, A. (2021): Exciting, Useful, Worrying, Futuristic: *Public Perception of Artificial Intelligence in 8 Countries*. 627–637. <https://doi.org/10.1145/3461702.3462605>
- Krägeloh, C. U. – Melekhov, V. – Alyami, M. M. – Medvedev, O. N. (2025): Artificial Intelligence Attitudes Inventory (AIAI): Development and validation using Rasch methodology. *Current Psychology*, 44(12), 12315–12327. <https://doi.org/10.1007/s12144-025-08009-1>
- Maier, S. B. – Jussupow, E. – Heinzl, A. (2020): Good, bad, or both? Measurement of physician’s ambivalent attitudes towards AI. 27th European Conference on Information Systems - Information Systems for a Sharing Society, *ECIS 2019*. <https://www.scopus.com/inward/record.uri?eid=2-s2.0-85087106822&partnerID=40&md5=5739b09a1755e70dd2160a7ab2a5532d>
- Maslej, N. – Fattorini, L. – Perrault, R. – Gil, Y. – Parli, V. – Kariuki, N. – Capstick, E. – Reuel, A. – Brynjolfsson, E. – Etchemendy, J. – Ligett, K. – Lyons, T. – Manyika, J. – Niebles, J. C. – Shoham, Y. – Wald, R. – Walsh, T. – Hamrah, A. – Santarlasci, L. – Lotufo, J. B. – Rome, A. – Shi, A. – Oak, S. (2025): Artificial Intelligence Index Report 2025 (Verzió 3). *arXiv*. <https://doi.org/10.48550/ARXIV.2504.07139>
- McGuire, W. J. (1985): *The nature of attitude and attitude change*. In: Lindzey, G. – Aronson, E. (eds.): *Handbook of Social Psychology* (Köt. 1–3., 223–349. p.). Random House.
- Mutanga, M. B. – Jugoo, V. – Adefemi, K. O. (2024): Lecturers’ Perceptions on the Integration of Artificial Intelligence Tools into Teaching Practice. *Trends in Higher Education*, 3(4), 1121–1133. <https://doi.org/10.3390/higheredu3040066>
- Nunnally, J. C. (1994): *Psychometric Theory 3E*. Tata McGraw-Hill Education. https://books.google.hu/books?id=6R_f3G58JsC
- OECD (2019): *Artificial Intelligence in Society*. OECD Publishing. <https://doi.org/10.1787/eed-fee77-en>
- Osgood, C. E. – Suci, G. J. – Tannenbaum, P. H. (1957): *The Measurement of Meaning*. University of Illinois Press. <https://books.google.hu/books?id=qk5qAAAAMAAJ>
- Park, J. – Woo, S. E. – Kim, J. (2024): Attitudes towards artificial intelligence at work: Scale development and validation. *Journal of Occupational and Organizational Psychology*, 97(3). <https://doi.org/10.1111/joop.12502>
- Pedraja-Rejas, L. – Lazo Vega, P. – Rojas Huanca, P. (2026): Navigating Ambivalence: Artificial Intelligence and Its Impact on Student Engagement in Engineering Education. *Behavioral Sciences*, 16(1). <https://doi.org/10.3390/bs16010011>
- Pinazo-Hernandis, S. – Carcavilla-González, N. (2026): Are future social workers ready for AI? Fears, barriers, and learning needs in higher education. *Social Work Education*. <https://doi.org/10.1080/02615479.2026.2631708>
- Rosenberg, M. J. – Hovland, C. I. (1960): *Cognitive, affective, and behavioral components of attitudes*. In: *Attitude Organization and Change: An Analysis of Consistency among Attitude Components* (1–14 p.). Yale University Press.
- Schepman, A. – Rodway, P. (2020): Initial validation of the general attitudes towards Artificial Intelligence Scale. *Computers in Human Behavior Reports*, 1, 100014. <https://doi.org/10.1016/j.chbr.2020.100014>

- 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). <https://doi.org/10.1080/10447318.2022.2085400>
- Schepman, A. – Rodway, P. (2026): Validation of the Short General Attitudes Towards Artificial Intelligence Scale: The Short GAAIS-10. *International Journal of Human-Computer Interaction*, 1–17. <https://doi.org/10.1080/10447318.2025.2610446>
- Schneider, I. K. – Novin, S. – van Harreveld, F. – Genschow, O. (2021): Benefits of being ambivalent: The relationship between trait ambivalence and attribution biases. *British Journal of Social Psychology*, 60(2). <https://doi.org/10.1111/bjso.12417>
- 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>
- Stein, J.-P. – Messingschlager, T. – Gnambs, T. – Hutmacher, F. – Appel, M. (2024): Attitudes towards AI: measurement and associations with personality. *Scientific Reports*, 14(1). <https://doi.org/10.1038/s41598-024-53335-2>
- Suh, W. – Ahn, S. (2022): Development and Validation of a Scale Measuring Student Attitudes Toward Artificial Intelligence. *Sage Open*, 12(2), 21582440221100463. <https://doi.org/10.1177/21582440221100463>
- van Harreveld, F. – Rutjens, B. T. – Schneider, I. K. – Nohlen, H. U. – Keskinis, K. (2014): In doubt and disorderly: Ambivalence promotes compensatory perceptions of order. *Journal of Experimental Psychology: General*, 143(4). <https://doi.org/10.1037/a0036099>
- 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). <https://doi.org/10.1080/10494820.2019.1674887>
- Yildiz, T. (2023): Measurement of Attitude in Language Learning with AI (MALL:AI). *Participatory Educational Research*, 10(4). <https://doi.org/10.17275/per.23.62.10.4>
- Zanna, M. – Rempel, J. (2008): *Attitudes: A new look at an old concept. Attitudes: Their structure, function, and consequences.*