



Közzététel: 2026. július 29.

A tanulmány címe:

A mesterséges intelligencia alapú ügyfélszolgálati megoldások elfogadása a telekommunikációs szektorban – egy pilotkutatás tanulságai

Szerző:

GREGUS-VOJTER NOÉMI, a Szegedi Tudományegyetem PhD-hallgatója

E-mail: vojternoemi97@gmail.com

DOI: <https://doi.org/10.20311/stat2026.07.hu0635>

Az alábbi feltételek érvényesek minden, a Központi Statisztikai Hivatal (a továbbiakban: KSH) Statisztikai Szemle c. folyóiratában (a továbbiakban: Folyóirat) megjelenő tanulmányra. Felhasználó a tanulmány vagy annak részei felhasználásával egyidejűleg tudomásul veszi a jelen dokumentumban foglalt felhasználási feltételeket, és azokat magára nézve kötelezőnek fogadja el. Tudomásul veszi, hogy a jelen feltételek megszegéséből eredő valamennyi kárért felelősséggel tartozik.

1. A jogszabályi tartalom kivételével a tanulmányok a szerzői jogról szóló 1999. évi LXXVI. törvény (Szt.) szerint szerzői műnek minősülnek. A szerzői jog jogosultja a KSH.
2. A KSH földrajzi és időbeli korlátozás nélküli, nem kizárólagos, nem átdadható, térítésmentes felhasználási jogot biztosít a Felhasználó részére a tanulmány vonatkozásában.
3. A felhasználási jog keretében a Felhasználó jogosult a tanulmány:
 - a) oktatási és kutatási célú felhasználására (nyilvánosságra hozatalára és továbbítására a 4. pontban foglalt kivétellel) a Folyóirat és a szerző(k) feltüntetésével;
 - b) tartalmáról összefoglaló készítésére az írott és az elektronikus médiában a Folyóirat és a szerző(k) feltüntetésével;
 - c) részletének idézésére – az átvevő mű jellege és célja által indokolt terjedelemben és az eredetihez híven – a forrás, valamint az ott megjelölt szerző(k) megnevezésével.
4. A Felhasználó nem jogosult a tanulmány továbbértékesítésére, haszonszerzési célú felhasználására. Ez a korlátozás nem érinti a tanulmány felhasználásával előállított, de az Szt. szerint önálló szerzői műnek minősülő mű ilyen célú felhasználását.
5. A tanulmány átdolgozása, újra publikálása tilos.
6. A 3. a)–c) pontban foglaltak alapján a Folyóiratot és a szerző(ke)t az alábbiak szerint kell feltüntetni:

„*Forrás: Statisztikai Szemle c. folyóirat 104. évfolyam 7. számában megjelent, Gregus-Vojter Noémi által írt, A mesterséges intelligencia alapú ügyfélszolgálati megoldások elfogadása a telekommunikációs szektorban – egy pilotkutatás tanulságai* című cikk (link csatolása)”

7. A Folyóiratban megjelenő tanulmányok kutatói véleményeket tükröznek, amelyek nem feltétlenül esnek egybe a KSH vagy a szerzők által képviselt intézmények hivatalos álláspontjával.

Gregus-Vojter Noémi

A mesterséges intelligencia alapú ügyfélszolgálati megoldások elfogadása a telekommunikációs szektorban – egy pilotkutatás tanulságai

Acceptance of artificial intelligence-based customer service solutions in the telecommunications sector – insights from a pilot study

Gregus-Vojter Noémi, a Szegedi Tudományegyetem PhD-hallgatója
E-mail: vojternoemi97@gmail.com

A tanulmány a mesterséges intelligencia (MI) alapú ügyfélszolgálati megoldások elfogadását vizsgálja a hazai telekommunikációs szektorban. A kutatás célja annak feltárása, hogy a különböző kognitív és érzelmi tényezők miként befolyásolják az automatizált rendszerekkel kapcsolatos használati szándékot. Az alkalmazott elméleti modell a bizalmat állítja a középpontba, olyan központi közvetítő elemként, amely összekapcsolja a kezdeti elvárásokat a tényleges döntéshozatali folyamattal. Az empirikus elemzés egy 104 fős válaszadói körben végzett pilotkutatás adataira támaszkodik, az adatok feldolgozása pedig strukturális egyenletmodellezéssel történt. Az eredmények rávilágítanak, hogy a szórakoztató használati élmény és a társas környezet hatása jelentősen erősíti a bizalmat, miközben a technológia emberszerűsége önmagában nem növeli a hitelességét. A vizsgálat legfontosabb megállapítása szerint a használati szándékot elsősorban a pozitív érzelmi reakciók motiválják, ami igazolja, hogy az elfogadás folyamata ebben a környezetben inkább érzelmi, mintsem tisztán racionális alapú. A tanulmány zárásként gyakorlati javaslatokat fogalmaz meg a szolgáltatók számára az ügyfélélmény javítása és a hatékony hibrid kiszolgálási modellek kialakítása érdekében.

Kulcsszavak: mesterséges intelligencia, technológiaelfogadási modellek, *AIDUA*, telekommunikáció, *PLS-SEM*

The study examines the acceptance of artificial intelligence (AI)-based customer service solutions in the domestic telecommunications sector. The aim of the research is to explore how different cognitive and emotional factors influence the intention to use automated systems. The applied theoretical model places trust at its core as a central mediating element that links initial expectations to the actual decision-making process. The empirical analysis is based on data from a pilot study conducted with 104 respondents, and the data were processed using structural equation modeling. The results highlight that the entertaining user experience and the influence of the social environment significantly enhance trust, while the anthropomorphism of the technology alone does not increase its credibility. The main finding of the study is that usage intention is primarily driven by positive emotional responses, confirming that the acceptance process in this context is more emotional than purely rational. Finally, the study formulates practical recommendations for service providers to improve customer experience and to develop effective hybrid service models.

Keywords: artificial intelligence, technology acceptance models, *AIDUA*, telecommunications, *PLS-SEM*

A mesterséges intelligenciával (MI) működő technológiák – mint például a szolgáltató robotok, a chatbotok és az autonóm döntéstámogató rendszerek – fejlesztése és széles körű alkalmazása az elmúlt évtizedekben jelentős mértékben felgyorsult, különösen az ügyfélkiszolgálási és szolgáltatásmenedzsment folyamatokban (Gursoy et al., 2019; Lin et al., 2020; Peng et al., 2022). Az MI-alapú megoldások integrálása a szolgáltatási láncba új minőségi dimenziókat nyitott meg: a gyorsabb és konzisztensebb ügyfél-interakciók, a nagy volumenű adatok valós idejű feldolgozása, valamint a személyre szabott szolgáltatásnyújtás lehetősége egyaránt hozzájárulnak a szervezeti hatékonyság és az ügyfélélmény javításához (Gursoy et al., 2019; Lin et al., 2020). A gyakorlatban számos iparágban megfigyelhető ezen technológiák térnyerése: a szállodaiparban például a Hilton Hotels a „Connie” nevű robotkonzultánst alkalmazza a vendéginterakciók támogatására, míg a kiskereskedelemben és az elektronikus kereskedelemben MI-alapú ajánlórendszerek és virtuális asszisztensek segítik a vásárlói döntéshozatalt (Gursoy et al., 2019). Ezeket a folyamatokat a szakirodalom gyakran a *smart service* (okos szolgáltatás) vagy *smart hospitality* (okos vendéglátás) koncepcióival írja le, ahol az intelligens technológiák a versenyelőny és az innováció meghatározó tényezőivé válnak (Wong et al., 2023).

Az MI-alapú rendszerek megvalósítása ugyanakkor nem tekinthető kizárólag technológiai vagy szervezeti kérdésnek, mivel sikerük alapvetően függ a felhasználók elfogadási hajlandóságától és tényleges használatától (Gursoy et al., 2019; Chi et al., 2023). Ennek megfelelően a technológiaelfogadás kutatása az innovációk terjedésének és alkalmazásának megértésére irányul, különös tekintettel arra, hogy a felhasználók milyen tényezők mentén döntenek az új technológiák elfogadásáról vagy elutasításáról (Keszey–Zsukk, 2017). A korábbi kutatások gyakran támaszkodtak hagyományos, funkcionális technológiai elfogadási modellekre, mint a *technology acceptance model* (TAM) (Davis, 1986) vagy az *unified theory of acceptance and use of technology* (UTAUT) (Venkatesh et al., 2003), azonban a szakirodalom egyetért abban, hogy ezek a keretrendszerek elégtelenek az MI elfogadásának átfogó magyarázatára (Gursoy et al., 2019; Chi et al., 2023; Manzoor et al., 2024). Ezen hiányosságok áthidalására Gursoy és szerzőtársai (2019) kidolgozták az AIDUA- (*artificially intelligent device use acceptance*) modellt, amely a kognitív értékelés elméletére (*cognitive appraisal theory*) alapozva egy többlépcsős, érzelmi alapú fogyasztói döntéshozatali folyamatot ír le (Gursoy et al., 2019; Wong et al., 2023; Chi et al., 2023). Az AIDUA-keretrendszer hangsúlyozza, hogy a vásárlók nem kizárólag racionális (kognitív) megfontolások alapján döntenek az MI elfogadásáról, hanem az érzelmi reakciók játsszák a kritikus, közvetítő szerepet a viselkedési szándék kialakításában (Manzoor et al., 2024).

Mindezek tükrében jelen kutatás a telekommunikációs szektor fogyasztói körében vizsgálja az MI-alapú ügyfélszolgálati rendszerek elfogadását. A kutatás

célja annak feltárása, hogy az AIDUA-modell konstrukciói miként befolyásolják a fogyasztók MI-alapú ügyfélszolgálati megoldásokkal kapcsolatos elfogadási hajlandóságát és használati szándékát. A vizsgálat empirikus pilotkutatás keretében valósult meg, amely során az eredeti AIDUA-modell egy kiterjesztett verziójának tesztelésére került sor online kérdőíves adatfelvétel alkalmazásával. Az adatelemzés a részleges legkisebb négyzetek strukturális egyenletmodellezésének módszerével (*partial least squares structural equation modeling*, PLS–SEM) történt, a SmartPLS szoftver alkalmazásával.

1. Szakirodalmi áttekintés

1.1. Az MI-alapú technológiák alkalmazási lehetőségei az ügyfélkiszolgálás területén

Az ügyfélkiszolgálás a szolgáltatási szektor egyik legkritikusabb érintkezési pontja, amely közvetlen hatást gyakorol az ügyfél-elégedettségre, a lojalításra és végső soron a vállalat versenyképességére (*Lacity–Willcocks, 2016*). A digitális transzformáció előrehaladtával a fogyasztók elvárásai jelentősen átalakultak: a gyors, rugalmas, személyre szabott és folyamatosan elérhető szolgáltatások iránti igény mára alapkövetelménnyé vált (*Lee–Lee, 2020*). Ebben a kontextusban az MI nem pusztán támogató technológiaként jelenik meg, hanem az ügyfélkiszolgálási folyamatok strukturális átalakulásának egyik meghatározó hajtóerejévé vált (*Wirtz et al., 2018*).

Az MI-alapú ügyfélkiszolgálási rendszerek legfontosabb előnye az automatizáció révén megvalósuló hatékony erőforrás-felhasználás. A természetes nyelvi feldolgozásra (*natural language processing*, NLP), illetve a gépi tanulásra épülő chatbotok és virtuális asszisztensek képesek nagyszámú ügyfélmegkeresés egyidejű kezelésére, jelentősen csökkentve a várakozási időt és az emberi erőforrásokra nehezedő terhelést (*Sharma, 2021; Tisóczy, 2022*). Ezek a rendszerek nem csupán előre definiált válaszokat adnak, hanem folyamatos tanulás révén egyre pontosabban képesek értelmezni a felhasználói szándékokat, kontextusfüggő válaszokat nyújtani, valamint összetettebb problémák megoldását is támogatni (*Duarte et al., 2022; Balmer et al., 2020*).

Az interaktív hangválaszrendszerek (*interactive voice response*, IVR) fejlődése jól szemlélteti az MI ügyfélkiszolgálásban betöltött szerepének átalakulását. Míg a korai IVR-megoldások elsősorban egyszerű menürendszerekre és billentyűalapú

navigációra épültek, addig a modern, MI-vel támogatott IVR-rendszerek már fejlett beszédfelismerési és nyelvértelmezési képességekkel rendelkeznek (*Corkrey–Parkinson, 2002; Wang et al., 2023*). Ennek eredményeként az ügyfelek természetes nyelven kommunikálhatnak a rendszerrel. Emellett képes azonosítani a problémát, releváns információkat szolgáltat, illetve szükség esetén az ügyintézés emberi operátorhoz továbbítja. Ez a hibrid megközelítés hozzájárul az ügyfélélmény javításához, miközben megőrzi az automatizáció hatékonysági előnyeit.

Az MI-alapú technológiák alkalmazása az ügyfélkiszolgálásban nem korlátozódik az interakciók automatizálására. A nagyméretű adathalmazok elemzésére képes rendszerek lehetővé teszik az ügyfélviselkedés mélyebb megértését, beleértve a preferenciák, szokások és panaszmintázatok azonosítását. Az így nyert információk alapján a vállalatok személyre szabott ajánlatokat, proaktív ügyfélmegkereséseket és prediktív támogatási megoldásokat alakíthatnak ki. Például az MI képes előre jelezni egy potenciális szolgáltatási problémát (*Santosh et al., 2020*) vagy ügyfél-elégedetlenséget, lehetőséget teremtve a megelőző beavatkozásra.

További jelentős alkalmazási területet jelentenek az ügyfél-azonosítást és a biztonságot támogató MI-megoldások. Az arcfelismerésen, a hangazonosításon vagy a viselkedésalapú biometrikus elemzésen alapuló rendszerek növelik az ügyfél-azonosítás pontosságát, miközben csökkentik a csalás és a visszaélések kockázatát (*Ouyang et al., 2021; Umamaheswari–Valarmathi, 2023*). Ezek a megoldások különösen nagy jelentőséggel bírnak olyan szektorokban, ahol érzékeny adatok kezelése történik, és az ügyfelek bizalmának fenntartása kulcsfontosságú.

Az MI ügyfélkiszolgálásban való alkalmazása ugyanakkor nem mentes a kihívásoktól. A technológiai korlátok mellett – mint a természetes nyelv komplexitása vagy a kontextusérzékeny kommunikáció nehézsége – kiemelt figyelmet kell fordítani az etikai és adatvédelmi szempontokra is. Az algoritmikus döntéshozatal átláthatósága, az adatkezelés jogszerűsége és a torzításmentes működés biztosítása elengedhetetlen feltétele az MI-alapú ügyfélkiszolgálási rendszerek hosszú távú elfogadásának (*Necz, 2023; Mittelstadt, 2019*).

Az MI ügyfélkiszolgálásban betöltött szerepe különösen markánsan jelenik meg a telekommunikációs szektorban, amelyet a magas ügyfélbázis, a komplex szolgáltatási portfólió, valamint az intenzív és gyakran problémacentrikus ügyfélinterakciók jellemeznek (*Pininti, 2024*). A távközlési szolgáltatók ügyfélszolgálatai naponta nagyszámú megkeresést kezelnek, amelyek egyaránt vonatkozhatnak technikai hibára, számlázási kérdésekre, szerződésmódosításokra vagy akár új szolgáltatások igénybevételére. Ez a heterogén igénystruktúra és az ügyintézés időkritikus jellege különösen alkalmassá teheti a szektort az MI-alapú automatizált megoldások alkalmazására.

A telekommunikációs ügyfélkiszolgálás egyik legfontosabb sajátossága a szolgáltatások technikai összetettsége. Az ügyfelek által jelzett problémák gyakran

nem egyszerű információhiányból, hanem hálózati, eszköz- vagy konfigurációs rendellenességekből fakadnak. Az MI-alapú rendszerek – különösen a tudásalapú chatbotok és döntéstámogató algoritmusok – képesek nagymennyiségű műszaki adat, hibajegy és korábbi megoldási minta elemzésére, ezáltal gyors és konzisztens válaszokat biztosítva. A prediktív analitikai megoldások révén a szolgáltatók akár előre is jelezhetik bizonyos hálózati problémák kialakulását, lehetőséget teremtve a proaktív ügyfélértesítésre és a megelőző beavatkozásra.

A szektor további jellegzetessége az ügyfélkapcsolatok hosszú távú jellege, amelyet gyakran szerződéses kötelezettségek és magas váltási költségek határoznak meg (Kim *et al.*, 2004). A váltási költségek ebben az esetben nem csupán pénzügyi terheket jelentenek, hanem magukban foglalják a szolgáltatóváltással járó adminisztratív folyamatokat, az esetleges hűségidőből fakadó kötbéreket, a technikai átállás nehézségeit, valamint az új szolgáltatások és rendszerek megszokásának időigényét is. Ennek következtében az ügyfélélmény minősége – különösen problémakezelési helyzetekben – kiemelt jelentőséggel bír az ügyfélmegtartás szempontjából. Ebben a kontextusban az érzelmek vizsgálata és értelmezése kulcsfontosságú tényezővé válik, mivel az ügyfél elégedettségét és lojalitását nem csupán a funkcionális megoldások, hanem az interakciók során átélt érzelmi tapasztalatok is jelentősen befolyásolják.

Az MI-alapú ügyfélszolgálati rendszerek alkalmazása ebben az összefüggésben kettős hatással bír: egyrészt javítja a válaszadási időt és az elérhetőséget (Satheesh–Nagaraj, 2021), másrészt azonban az automatizált interakciók személytelensége negatív érzelmi reakciókat is kiválthat. Ez tovább erősíti annak szükségességét, hogy az ügyfél-interakciók során az érzelmi állapotok felismerése és kezelése beépüljön a rendszertervezésbe. Különösen igaz ez olyan helyzetekben, ahol az ügyfelek empátiát, megértést vagy egyedi mérlegelést várnak el, ami rávilágít az MI és az emberi ügyintézők együttműködésének fontosságára (Lu *et al.*, 2020).

A telekommunikációs szektorban alkalmazott MI-alapú ügyfélszolgálati megoldások elfogadását továbbá jelentősen befolyásolják az adatvédelemmel és adatbiztonsággal kapcsolatos aggályok. A szolgáltatók nagy mennyiségű személyes és forgalmi adatot kezelnek, így az MI-rendszerek működésével kapcsolatos átláthatóság, valamint az algoritmikus döntéshozatalba vetett bizalom kulcsfontosságú tényezővé válik. Ebben a kontextusban a bizalom nem csupán kiegészítő tényező, hanem alapvető előfeltétele az MI-alapú rendszerek elfogadásának, mivel meghatározza, hogy az ügyfelek mennyire tartják megbízhatónak, biztonságosnak és legitimnek az automatizált ügyintézést. Az ügyfelek elfogadási hajlandóságát ezért jelentősen csökkentheti, ha az MI-alapú rendszer működése „fekete dobozként” jelenik meg, vagy ha nem egyértelmű, hogy egy adott döntést automatizált algoritmus vagy emberi szereplő hozott meg.

A telekommunikációs ügyfélszolgálatok esetében ezért egyre inkább előtérbe kerülnek a hibrid szolgáltatási modellek, amelyekben az MI-alapú rendszerek elsősorban az egyszerű, rutinszerű ügyek kezelését végzik, míg a komplexebb, érzelmileg megterhelőbb esetek emberi ügyintézőkhöz kerülnek továbbításra. Ez a megközelítés nemcsak az operatív hatékonyságot növeli, hanem hozzájárulhat az ügyfelek technológiába vetett bizalmának erősítéséhez is. E sajátosságok alapján a telekommunikációs szektor különösen alkalmas kutatási környezetet biztosít az MI-alapú ügyfélszolgálati rendszerek elfogadásának vizsgálatára, mivel egyszerre jelennek meg benne a technológiai hatékonysági elvárások, az adatbiztonsági kockázatok és az érzelmi alapú fogyasztói reakciók. Ennek megfelelően a következő fejezet az MI-alapú rendszerek fogyasztói elfogadásának elméleti kereteit és az alkalmazott modelleket mutatja be.

1.2. Az MI-alapú ügyfélszolgálati rendszerek fogyasztói elfogadása

Az automatizált, adaptív és gyakran emberszerű interakciókra képes rendszerek nem csupán az operatív hatékonyságot növelik, hanem jelentősen átalakítják a szolgáltatás észlelt minőségét és a fogyasztói élményt is. Ezen technológiák sikeres alkalmazásának ugyanakkor előfeltétele annak megértése, hogy milyen pszichológiai és kognitív mechanizmusok vezetnek a fogyasztói elfogadáshoz vagy éppen elutasításhoz.

A kutatások túlnyomó része olyan hagyományos keretrendszerekre támaszkodott, mint a TAM vagy az UTAUT (Gursoy et al., 2019; Chi et al., 2023). Bár ezen modellek évtizedeken keresztül meghatározó szerepet töltek be az információs rendszerek kutatásában, az MI-alapú rendszerek esetében több szempontból is korlátozott magyarázóerővel rendelkeznek (Gursoy et al., 2019; Della Corte et al., 2023).

Egyrészt ezeket a keretrendszereket alapvetően nem intelligens technológiák – például mobil bejelentkezési rendszerek, érintésmentes fizetési megoldások vagy hagyományos információs rendszerek – vizsgálatára hozták létre (Gursoy et al., 2019; Chow et al., 2023). Ezzel szemben az MI-alapú megoldások adaptív, tanulásra képesek és gyakran emberre jellemző viselkedést mutatnak, amik alapjaiban változtatják meg az ember–technológia interakció jellegét (Gursoy et al., 2019; Chow et al., 2023; Roy et al., 2024). Másrészt a TAM és az UTAUT központi konstrukciói – az észlelt könnyű használhatóság, illetve az erőfeszítés-elvárás – elsősorban a technológia működtetésének elsajátítására fókuszálnak (Gursoy et al., 2019; Chow et al., 2023). Az MI-alapú rendszerek, például a természetes nyelven kommunikáló chatbotok, gyakran nem igényelnek explicit tanulási folyamatot a

felhasználók részéről, így ezen változók jelentősége részben háttérbe szorul (Gursoy et al., 2019; Roy et al., 2024).

További lényeges eltérés, hogy a fogyasztók az MI teljesítményét nem egy statikus technológiai eszközhöz, hanem az ügyfélszolgálati munkatársak által nyújtott szolgáltatási szinthez viszonyítják (Gursoy et al., 2019; Roy et al., 2024).

Végül a klasszikus elfogadási modellek nagyrészt figyelmen kívül hagyják az interakció társadalmi és érzelmi beágyazottságát, valamint az olyan tényezőket, mint az antropomorfizmus (emberszerűség), amelyek az MI-rendszerekkel való kapcsolatban kiemelt szerepet játszanak (Gursoy et al., 2019; Della Corte et al., 2023).

Többek között ezen hiányosságok áthidalására dolgozták ki Gursoy és szerzőtársai (2019) az AIDUA-modellt. A modell a kognitív értékelés elméletére (*cognitive appraisal theory*) támaszkodva (Gursoy et al., 2019; Roy et al., 2024; Begum et al., 2025) egy háromlépcsős folyamatot ír le az elfogadás generálására: elsődleges, valamint másodlagos értékelés és az eredményfázis. Ez a modell megerősíti, hogy a döntéshozatal nem egy pillanatnyi reakció, hanem egy összetett kognitív értékelési mechanizmusból fakad (Gursoy et al., 2019; Roy et al., 2024; Manzoor et al., 2024).

2. A kutatás konceptuális modellje és hipotézisei

2.1. A kutatásban alkalmazott módosított modell

A kutatás elméleti keretrendszere egy többszintű modellfejlődési folyamatra épül, amelynek kiindulópontja a Gursoy és szerzőtársai (2019) által kidolgozott AIDUA-modell. Ez a keretrendszer a kognitív értékelés elméletére alapozva egy háromlépcsős folyamatként írja le az MI-elfogadást:

1. Elsődleges értékelés (*primary appraisal*)

Az elsődleges értékelési szakaszban az ügyfelek három fő tényező alapján mérik fel az MI-eszközök használatának fontosságát és relevanciáját: társadalmi hatás, a hedonikus motiváció és az antropomorfizmus (Gursoy et al., 2019; Roy et al., 2024; Begum et al., 2025; Manzoor et al., 2024).

- a) Társadalmi hatás (*social influence*): A szakirodalmi forrásokban következetesen láthatjuk, hogy a társadalmi hatás pozitívan kapcsolódik a teljesítményelváráshoz (Gursoy et al., 2019; Della Corte et al., 2023). Ha a társadalmi hálózat (például barátok, család) kedvező véleményt nyilvánít az MI-ről, az ügyfél eleve pozitívabb teljesítményt vár el (Gursoy

et al., 2019). A társadalmi hatás és az erőfeszítés-elvárás közötti kapcsolat azonban nagyrészt nem szignifikáns. Ez arra utal, hogy a felhasználók inkább a potenciális előnyökre, mint a költségekre összpontosítanak, ha a viselkedés megfelel a csoportnormáknak (*Gursoy et al.*, 2019).

- b) Hedonikus motiváció (*hedonic motivation*): A használatból várt szórakozást és élvezetet tükrözi, a teljesítmény- és erőfeszítés-elvárás legerősebb előrejelzője. Pozitívan kapcsolódik a teljesítményelváráshoz, mivel a felhasználók, akik élvezetesnek találják az MI-t, hajlamosak a használatát előnyben részesítő értékelést végezni (*Gursoy et al.*, 2019). Ugyanakkor a hedonikus motiváció negatívan kapcsolódik az erőfeszítés-elváráshoz, ami azt jelenti, hogy a motivált felhasználók kevésbé érzékelik a technológia használatát nehézkesnek (*Wong et al.*, 2023).
 - c) Antropomorfizmus (*anthropomorphism*): Az MI humánszerű jellemzőire (megjelenés, tudat, érzelem) utal. Az antropomorfizmus pozitívan kapcsolódik az erőfeszítés-elváráshoz. A humán jellegű vonások felerősítik a szükséges erőfeszítés észlelt mennyiségét, mivel a felhasználónak egyszerre kell tanulnia egy technológiai eszköz használatát, és emberi lényként kell kezelnie az MI-t (*Gursoy et al.*, 2019). Az antropomorfizmus és a teljesítményelvárás között azonban az eredeti modellben nem volt szignifikáns kapcsolat. Ez arra utal, hogy a humánszerű vonások nem feltétlenül járnak előnnyel a szolgáltatás teljesítménye vagy a várható szolgáltatás minősége szempontjából. A humán jellegű jegyek nem növelik a teljesítményelvárást, vagyis az ügyfelek nem várnak el jobb feladatvégzést attól, mert az MI emberibbnek tűnik (*Gursoy et al.*, 2019). Ez gyakorlati szempontból azt a kockázatot hordozhatja a szolgáltatást tervező menedzserek számára, hogy az MI túlzott antropomorfizálása növelheti a felhasználói kognitív terhelést és az észlelt használati erőfeszítést anélkül, hogy egyidejűleg javítaná a szolgáltatás teljesítményével vagy hatékonyságával kapcsolatos elvárásokat. Azonban ez a kapcsolat az egyik leginkább ellentmondásos pont a vizsgált szakirodalomban: miközben többen megerősítették ezt a megállapítást, voltak, akik pozitív összefüggést tártak fel az antropomorfizmus és a teljesítményelvárás között (*Wong et al.*, 2023). Ezek az eltérések azt jelezhetik, hogy az MI emberszerű tulajdonságainak megítélése nagymértékben függ a kulturális és kontextuális környezettől.
2. Másodlagos értékelés (*secondary appraisal*)
- A másodlagos értékelés során a felhasználók a teljesítményelvárás és az erőfeszítés-elvárás alapján végzik el az MI-használat szisztematikus értékelését, ami pozitív vagy negatív érzelmek kialakulásához vezet (*Gursoy et al.*, 2019; *Manzoor et al.*, 2024).

- a) Teljesítményelvárás (*performance expectancy*): Pozitívan hat az érzelmekre. A gyors, pontos és következetes szolgáltatás ígérete pozitív érzelmeket generál (Gursoy et al., 2019; Li et al., 2024).
- b) Erőfeszítés-elvárás (*effort expectancy*): Negatívan hat az érzelmekre. A túl sok energiát igénylő interakció negatív érzelmekhez vezethet (Gursoy et al., 2019; Li et al., 2024).
- c) Érzelmek (*emotion*): A másodlagos értékelés elengedhetetlen lezárását jelentik, mielőtt a fogyasztó meghozná a végső viselkedési döntését (Gursoy et al., 2019).

A teljesítményelvárás hatása az érzelmekre azonban lényegesen erősebb, mint az erőfeszítés-elvárás negatív hatása. Ez azt jelentheti, hogy az ügyfelek számára a szolgáltatás teljesítménye fontosabb motivációs tényező, mint a tapasztalt erőfeszítés (Gursoy et al., 2019).

3. Eredményfázis (*outcome stage*)

A végső kimenet az MI-eszközök iránti érzelmi állapotban gyökerezik, ami meghatározza a hajlandóságot a használatra, vagy a kifogást a használattal szemben.

- a) Hajlandóság az MI-eszközök használatának elfogadására (*willingness to accept the use of AI devices*): A pozitív érzelmek pozitívan befolyásolják a hajlandóságot az elfogadásra (Gursoy et al., 2019; Manzoor et al., 2024).
- b) Kifogás az MI-eszközök használatával szemben (*objection to use of AI devices*): A domináló negatív érzelmekkel (például frusztráció, félelem, bizonytalanság, szorongás vagy aggodalom) rendelkező ügyfeleknél nagyobb valószínűséggel merül fel kifogás az MI-vel szemben (Manzoor et al., 2024).

Az eredeti modell egyik legfontosabb megállapítása, hogy az MI elfogadása során az érzelmi reakciók kritikus közvetítő szerepet töltenek be a kognitív folyamatok és a tényleges használati szándék között.

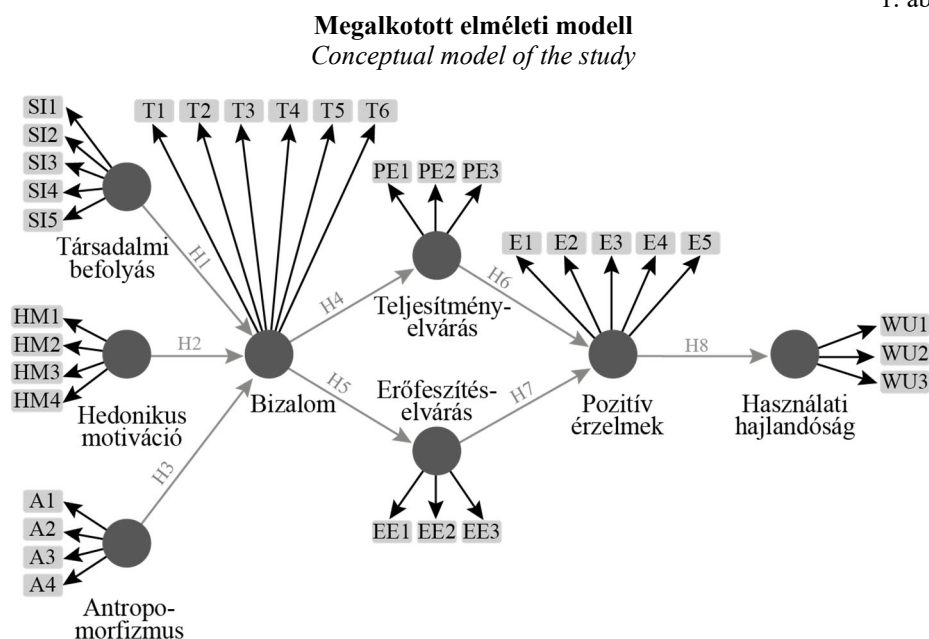
Mivel az alapmodell nem tartalmazta explicit módon a bizalom tényezőjét, a szakirodalmi fejlődés következő lépcsőfokát a Della Corte és szerzőtársai (2023) által kiterjesztett AIDUA-modell jelentette. Ez a megközelítés a bizalmat a fogyasztói döntéshozatal központi mediátoraként integrálta a rendszerbe, amely hídként kapcsolja össze a heurisztikus tényezőket (például az antropomorfizmust vagy a hedonikus motivációt) a kognitív értékelési folyamatokkal. A bizalom integrálása különösen kritikus a telekommunikációs szektorban, ahol a szolgáltatók hatalmas mennyiségű szenzitív adatot kezelnek, beleértve a személyes ügyféladatakat és a hálózati információkat, így az adatbiztonság és a megbízható működés alapvető elvárás (Sen et al., 2024).

Jelen tanulmányban a *Della Corte és szerzőtársai (2023)* által javasolt kiterjesztett modellt módosítom tovább. A módosítás során elhagytam az egyéni szintű változókat – mint az észlelt kockázatot és az innovativitást –, hogy a vizsgálat fókuszba célzottan a bizalom mediáló szerepére, valamint a kognitív és érzelmi útvonalak hatásmechanizmusára irányuljon a telekommunikációs szektor ügyfélkiszolgálási kontextusában. Ez a finomított struktúra lehetővé teszi a modell stabilitásának és értelmezhetőségének pontosabb tesztelését a választott empirikus környezetben.

Így az általam alkalmazott kiterjesztett modell a következő struktúrát követi:

1. Heurisztikus tényezők: A modell kiindulópontját a korábban bemutatott AIDUA-keretrendszerből származó heurisztikus tényezők alkotják, amelyek a társadalmi befolyást, a hedonikus motivációt és az antropomorfizmust foglalják magukban. Ezek a változók az elsődleges kognitív értékelés alapját képezik, és meghatározzák a felhasználók kezdeti észlelését az MI-alapú ügyfélszolgálati rendszerekkel kapcsolatban.
2. A bizalom mediáló szerepe: A bizalom a modell központi mediáló konstrukciójaként jelenik meg, amely összeköti a heurisztikus tényezőket a kognitív értékelési folyamatokkal. A magasabb szintű bizalom növeli a teljesítményelvárást, miközben csökkenti az erőfeszítés-elvárást, ezáltal kedvezőbb pszichológiai értékelési keretet biztosít a technológia használatához (*Della Corte et al., 2023*).
3. Kognitív értékelési tényezők: A kognitív értékelési szakaszban a teljesítményelvárás és az erőfeszítés-elvárás határozza meg a rendszerrel kapcsolatos észlelt hasznosságot és használati nehézséget. Ezek a tényezők közvetlenül befolyásolják az érzelmi reakciók kialakulását, és a bizalom hatásmechanizmusán keresztül kapcsolódnak a heurisztikus előzményekhez.
4. Érzelmi kimenet: A kognitív értékelés eredményeként kialakuló érzelmi reakciók közvetítő szerepet töltenek be a modellben. A pozitív érzelmek elősegítik a technológia elfogadását, míg a negatív érzelmek csökkentik annak elfogadási valószínűségét (*Della Corte et al., 2023*).
5. Használati szándék: A modell végső kimeneti változója a használati szándék, amely az érzelmi reakciók közvetlen következményeként jelenik meg, és az MI-alapú ügyfélszolgálati rendszerek elfogadásának viselkedési indikátorát jelenti.

1. ábra



Forrás: saját szerkesztés Della Corte és szerzőtársai (2023) alapján.

2.2. A kutatás hipotézisei

A modellben szereplő kapcsolatok vizsgálata érdekében a kutatás a következő hipotéziseket fogalmazza meg:

H1: A társadalmi befolyás pozitív hatással van az MI-alapú ügyfélszolgálati rendszerek iránti bizalomra.

A társas környezet véleménye jelentős szerepet játszik az új technológiák elfogadásában, különösen olyan helyzetekben, amikor a felhasználók korlátozott tapasztalattal rendelkeznek az adott rendszerrel kapcsolatban (*Gursoy et al., 2019; Roy et al., 2024*).

H2: A hedonikus motiváció pozitív hatással van az MI-alapú ügyfélszolgálati rendszerek iránti bizalomra.

Az MI használatához kapcsolódó élmény és szórakoztató jelleg kedvező attitűdöt alakíthat ki a technológiával szemben, ami hozzájárulhat a bizalom kialakulásához (*Wong et al., 2023; Ribeiro et al., 2022*).

H3: Az antropomorfizmus pozitív hatással van az MI-alapú ügyfélszolgálati rendszerek iránti bizalomra.

Az emberszerű tulajdonságok növelhetik a rendszer hitelességét és empatikus jellegének észlelését, ami elősegítheti a bizalom kialakulását (*Chow et al., 2023; Della Corte et al., 2023*).

H4: A bizalom pozitív hatással van a teljesítményelvárásra.

A technológiába vetett magasabb bizalmi szint növeli annak valószínűségét, hogy a felhasználók hatékony és megbízható szolgáltatást várjanak az MI-alapú rendszertől (*Ribeiro et al., 2022; Chi et al., 2023; Della Corte et al., 2023*).

H5: A bizalom negatív hatással van az erőfeszítés-elvárásra.

A bizalom csökkenti a technológiával kapcsolatos bizonytalanságot, aminek következtében a felhasználók kevésbé érzékelik bonyolultnak vagy nehézkesnek a rendszer használatát (*Ribeiro et al., 2022; Chi et al., 2023; Della Corte et al., 2023*).

H6: A teljesítményelvárás pozitív hatással van a pozitív érzelmek kialakulására.

A gyors, pontos és következetes szolgáltatás észlelése kedvező érzelmi reakciókat válthat ki a felhasználókban, ezáltal erősítve a pozitív érzelmek kialakulását (*Gursoy et al., 2019; Khan et al., 2023; Roy et al., 2024; Schiavo et al., 2024*).

H7: Az erőfeszítés-elvárás negatív hatással van a pozitív érzelmek kialakulására.

Amennyiben a rendszer használata túlzott mentális vagy technikai erőfeszítést igényel, az csökkentheti a felhasználók pozitív érzelmeit, illetve kedvezőtlen érzelmi reakciókat válthat ki (*Gursoy et al., 2019; Wong et al., 2023; Della Corte et al., 2023*).

H8: A pozitív érzelmek pozitív hatással vannak az MI-alapú ügyfélszolgálati rendszerek használati szándékára.

A kedvező érzelmi reakciók növelik a technológia elfogadásának és jövőbeli használatának valószínűségét (*Gursoy et al., 2019; Wong et al., 2023; Schiavo et al., 2024*).

A kutatás hipotéziseinek és a konceptuális modell kapcsolatainak empirikus tesztelése PLS–SEM módszertannal történt, amely lehetővé tette a látens konstrukciók közötti kapcsolatok egyidejű vizsgálatát, valamint pilot jellegű kutatások esetén is megfelelő elemzési keretet biztosít.

3. Kutatási módszertan

Jelen kutatás célja az MI-alapú ügyfélszolgálati rendszerek fogyasztói elfogadását leíró, *Della Corte és szerzőtársai (2023)* által kiterjesztett, de módosított AIDUA-modell empirikus tesztelése, egy pilotkutatás keretében. A vizsgálat kvantitatív

kutatási megközelítést alkalmaz, ami lehetővé teszi a modellben szereplő változók közötti összefüggések statisztikai elemzését.

3.1. Minta és adatgyűjtés

A kutatás keresztmetszeti (*cross-sectional*) adatfelvételen alapult, amely egy adott időpontban vizsgálja a válaszadók MI-alapú ügyfélszolgálati rendszerekkel kapcsolatos észleléseit és viselkedési szándékait. A keresztmetszeti dizájn lehetővé teszi, hogy a vizsgált konstrukciók közötti kapcsolatok egy időpillanatban kerüljenek elemzésre, anélkül, hogy longitudinális követésre sor kerülne.

Az empirikus vizsgálat online kérdőíves adatgyűjtéssel valósult meg, amelyet 2026 áprilisa és májusa között tettem elérhetővé a célcsoport számára. Az adatfelvétel a telekommunikációs szolgáltatásokat igénybe vevő fogyasztók körében történt. A válaszadók elérése közösségi média platformokon, azon belül három, Magyarország legnagyobb piaci részesedéssel (*NMHH, 2025*) rendelkező telekommunikációs szolgáltatóhoz kapcsolódó, nem hivatalos Facebook-csoporton keresztül zajlott: a Yettel Magyarország, a One Magyarország és a Magyar Telekom felhasználói közösségeiben. A csoportok összesített tagsága több tízezer fő, így a felület alkalmas volt olyan válaszadók elérésére, akik tényleges tapasztalattal rendelkeznek telekommunikációs ügyfélszolgálati rendszerekkel, beleértve a chatbotokat és egyéb automatizált ügyintézési megoldásokat.

A mintavétel Facebook-csoportokon keresztül, önkéntes részvételen alapuló módon történt, amelyből adódóan a minta nem tekinthető statisztikailag reprezentatívnak a teljes telekommunikációs ügyfél-populációra nézve. A kutatás ugyanakkor pilot jelleggel készült, és célja elsősorban a vizsgált modell empirikus tesztelése volt, amelyhez a rendelkezésre álló minta megfelelő alapot biztosít.

A kérdőívet 2026. április 3-án osztottam meg, a válaszok gyűjtése 2026. május 19-ig tartott. A kérdőív kitöltése önkéntes és anonim módon történt. Az adatfelvétel lezárásáig összesen 118 válasz érkezett, amelyek adattisztítási eljárásokon estek át. Az elemzésből kizártam 14 olyan kitöltő választ, akik nem rendelkeztek releváns, MI-alapú ügyfélszolgálati tapasztalattal, illetve nem ismerték az MI fogalmát. Emellett további egy választ távolítottam el az adathalmazból az egysíkú válaszadási mintázat (*straight-lining*) miatt. Az adat-előkészítés során a T4 jelölésű változó átkódolása is megtörtént a további statisztikai elemzések megbízhatóságának biztosítása érdekében. Az adattisztítást követően 104 válasz alkotta a mintát, ez képezte az összes további statisztikai elemzés alapját, ami ugyan nem tekinthető véletlen mintavételnek, ezért az eredményeket csak óvatosan lehet, ugyanakkor ez az elemszám már megfelelő bizonyos következtetések levonásához.

A kutatás pilot jellegéből adódóan a minta mérete korlátozott, azonban alkalmas a mérési modell és a strukturális kapcsolatok, azaz a külső és a belső modell előzetes tesztelésére (Kazár, 2014). A végső mintanagyság a PLS–SEM módszertani követelményei alapján került értékelésre.

A minta demográfiai megoszlása alapján a válaszadók nemek szerinti aránya kiegyensúlyozott, a férfiak enyhe túlsúlyával. A korosztály szerinti eloszlás azt mutatja, hogy a legtöbb válaszadó az 1965–1980 és az 1981–1996 közötti korcsoportba tartozott, míg az 1946–1964 közötti és az 1997–2012 közötti korcsoport kisebb arányban képviseltették magukat. A végzettségi adatok alapján a minta többsége felsőfokú végzettséggel rendelkezik, ami a vizsgált online csatornák jellegéből adódóan digitálisan aktív populációt jelez.

1. táblázat

A minta demográfiai jellemzői (N=104)
Demographic characteristics of the sample (N=104)

Férfiak (N = 68)	Alapfok	Középfok	Felsőfok	Egyéb	Összesen
1946–1964	0 (0,0%)	2 (1,9%)	1 (0,9%)	0 (0,0%)	3 (2,9%)
1965–1980	0 (0,0%)	10 (9,6%)	14 (13,5%)	1 (0,9%)	25 (24,0%)
1981–1996	0 (0,0%)	16 (15,4%)	9 (8,7%)	0 (0,0%)	25 (24,0%)
1997–2012	3 (2,9%)	4 (3,8%)	8 (7,7%)	0 (0,0%)	15 (14,4%)
Összesen	3 (2,9%)	32 (30,8%)	32 (30,8%)	1 (0,9%)	68 (65,4%)
Női (N = 36)	Alapfok	Középfok	Felsőfok	Egyéb	Összesen
1946–1964	0 (0,0%)	1 (0,9%)	1 (0,9%)	0 (0,0%)	2 (1,9%)
1965–1980	1 (0,9%)	3 (2,9%)	7 (6,7%)	0 (0,0%)	11 (10,6%)
1981–1996	0 (0,0%)	4 (3,8%)	7 (6,7%)	1 (0,9%)	12 (11,5%)
1997–2012	0 (0,0%)	2 (1,9%)	9 (8,7%)	0 (0,0%)	11 (10,6%)
Összesen	1 (0,9%)	10 (9,6%)	24 (23,1%)	1 (0,9%)	36 (34,6%)

Forrás: saját szerkesztés.

3.2. Adatelemzési eljárás

A kutatás során az adatok statisztikai elemzését a SmartPLS szoftver segítségével végeztem. Az empirikus vizsgálat középpontjában a módosított AIDUA-modell előzetes, pilot jellegű tesztelése állt, amelyhez PLS–SEM egyenletmodellezést alkalmaztam.

A PLS–SEM-módszer alkalmazását indokolta a kutatás feltáró jellege, a viszonylag alacsony mintanagyság (Haenlein–Kaplan, 2004), valamint a több látens konstrukcióból álló komplex modell vizsgálata. A módszer lehetővé teszi a belső

és a külső modell egyidejű elemzését, ezért különösen alkalmas újonnan adaptált vagy módosított elméleti modellek előzetes empirikus tesztelésére.

Az elemzés első szakaszában a belső modell értékelése történt meg. Ennek során vizsgáltam az indikátorok faktorterheléseit, a konstrukciók belső konzisztenciáját, valamint a konvergencia és diszkriminációs validitás teljesülését. A megbízhatóság ellenőrzésére Cronbach-alfa és kompozit megbízhatósági (*composite reliability*, CR) mutatók kerültek alkalmazásra, míg a konvergencia validitás vizsgálata az átlagos kivonatolt variancia (*average variance extracted*, AVE) alapján történt.

A második szakaszban a strukturális, vagyis a külső modell előzetes tesztelésére került sor, amely során a konstrukciók közötti kapcsolatok erősségét és irányát vizsgáltam. Az elemzés keretében értékeltem az útkoefficienseket (*path coefficients*), a modell magyarázóerejét (R^2), valamint a kapcsolatok statisztikai szignifikanciáját *bootstrap* eljárás segítségével. A hipotézisek vizsgálata ugyanakkor jelen kutatásban másodlagos szerepet töltött be; az elemzés elsődleges célja a módosított modell működésének, valamint az egyes konstrukciók közötti kapcsolatok előzetes empirikus feltérképezése volt.

3.3. Mérési módszertan, vagyis a belső modell

A kutatás során vizsgált látens konstrukciókat, az azokhoz tartozó mérési indikátorokat, valamint a hozzájuk rendelt kódokat a 2. táblázat szemlélteti. Az egyes mérési tételek kialakítása a releváns nemzetközi szakirodalomban alkalmazott, validált skálák alapján történt, amelyeket a kutatás telekommunikációs és MI-alapú ügyfélszolgálati kontextusához igazítottam. A konstrukciók operacionalizálása során kiemelt szempont volt, hogy az indikátorok alkalmasak legyenek a válaszadók MI-alapú ügyfélszolgálati rendszerekkel kapcsolatos attitűdjeinek, észleléseinek és viselkedési szándékainak megragadására.

A kérdőívben szereplő állításokat a válaszadók ötfokú Likert-skálán értékelték, ahol az „1 = egyáltalán nem értek egyet”, míg az „5 = teljes mértékben egyetértek” válaszlehetőséget jelentette. A pozitív érzelmek konstrukció esetében szemantikus differenciál skála került alkalmazásra, amely az érzelmi reakciók intenzitását és irányát mérte ellentétpárok segítségével.

Az empirikus elemzést megelőzően a mérési modell előzetes validálása is megtörtént. Ennek során vizsgáltam az egyes indikátorok faktorterheléseit, a konstrukciók belső konzisztenciáját, valamint a konvergencia és diszkriminációs validitás teljesülését. Az előzetes elemzések alapján azok az indikátorok kerültek volna eltávolításra a végső modellből, amelyek nem felelnek meg a PLS–SEM módszertani követelményeinek, például alacsony faktorterhelés vagy jelentős keresztterhelés esetén. A végső modellbe kizárólag a megfelelő megbízhatósági és validitási kritériumoknak megfelelő indikátorok kerültek bevonásra.

2. táblázat

Konstrukciók és mérési indikátorok
Constructs and measurement indicators

Konstrukció	Mérési indikátor	Kód
Antropomorfizmus (A)	Az MI-alapú ügyfélkiszolgáló rendszerek saját „gondolkodással” rendelkeznek.	A1
	Az MI-alapú ügyfélkiszolgáló rendszerek tudattal bírnak.	A2
	Az MI-alapú ügyfélkiszolgáló rendszerek saját döntési szabadsággal rendelkeznek.	A3
	Az MI-alapú ügyfélkiszolgáló rendszerek képesek érzelmeket átélni.	A4
Pozitív érzelmek (E)	Unott – Nyugodt	E1
	Késégsbeesett – Reményteljes	E2
	Bosszús – Elégedett / elragadtatott	E3
	Melankolikus / lehangolt – Derűs	E4
	Elégedetlen – Elégedett	E5
Erőfeszítés-elvárás (EE)	Az MI-alapú ügyfélkiszolgáló rendszerek használata túl sok időt vesz igénybe.	EE1
	Sok időbe telik megtanulni, hogyan lehet hatékonyan használni az MI-alapú ügyfélkiszolgáló rendszereket.	EE2
	Az MI-alapú ügyfélkiszolgáló rendszerekkel való interakció nehezen érthető és bonyolult.	EE3
Hedonikus motiváció (HM)	Az MI-alapú ügyfélkiszolgáló rendszerekkel való interakció szórakoztató.	HM1
	Az MI-alapú ügyfélkiszolgáló rendszerek használata élvezetes.	HM2
	Az MI-alapú ügyfélkiszolgáló rendszerekkel való interakció kellemes élményt nyújt.	HM3
	Maga az interakció folyamata kellemes lenne számomra.	HM4
Teljesítmény-elvárás (PE)	Antropomorfizmus ügyfélkiszolgáló rendszerek által nyújtott információk pontosabbak, mint az emberi ügyintézőké.	PE1
	Az MI-alapú ügyfélkiszolgáló rendszerek kevesebb emberi hibával szolgáltatnak információt.	PE2
	Az MI-alapú ügyfélkiszolgáló rendszerek következetesebb információkat nyújtanak, mint az emberi ügyintézők.	PE3
Bizalom (T)	Általánosságban bízok a mesterséges intelligenciában.	T1
	A mesterséges intelligencia számos probléma megoldásában segítséget nyújt.	T2
	Jó ötletnek tartom, ha a mesterséges intelligencia támogatására támaszkodok.	T3
	Nem bízok a mesterséges intelligencia által nyújtott információkban.	T4
	A mesterséges intelligenciát megbízhatónak tartom.	T5
	Hajlandó vagyok a mesterséges intelligenciára támaszkodni.	T6

(A táblázat folytatása a következő oldalon)

(folytatás)

Konstrukció	Mérési indikátor	Kód
Társadalmi befolyás (SI)	A számomra fontos emberek elvárnák, hogy használjam az MI-alapú ügyfélkiszolgáló rendszereket.	SI1
	A környezetemben az MI-alapú ügyfélkiszolgáló rendszereket használó emberek nagyobb presztízzsel rendelkeznek.	SI2
	Azok az emberek, akiknek a véleményét nagyra értékelem, támogatnák az MI-alapú ügyfélkiszolgáló rendszerek használatát.	SI3
	A számomra fontos személyek bátorítanának az MI-alapú ügyfélkiszolgáló rendszerek használatára.	SI4
	A társas környezetemben az MI-alapú ügyfélkiszolgáló rendszereket használók magasabb társadalmi státusszal rendelkeznek.	SI5
Használati hajlandóság (WU)	Hajlandó vagyok információt igénybe venni az MI-alapú ügyfélkiszolgáló rendszerektől.	WU1
	Az MI-alapú ügyfélkiszolgáló rendszerekkel való interakció során jól érzem magam.	WU2
	Az MI-alapú ügyfélkiszolgáló rendszerekkel való interakció során könnyen kizárom a környezetem.	WU3

Forrás: saját szerkesztés.

A mérési modell ellenőrzése során megvizsgáltam az indikátorok multikollinearitását és megbízhatóságát. Az elemzés eredménye alapján az E3 indikátor törésre került, mivel a hozzá tartozó VIF-érték meghaladta az elfogadható küszöbértéket (6,149), ami túlzott multikollinearitásra és redundáns információtartalomra utalt. Az indikátor eltávolításával javult a modell illeszkedése és a konstrukció mérési tisztasága.

4. Eredmények

4.1. A mérési modell értékelése

A kutatás során első lépésként a mérési modell megbízhatóságának és validitásának vizsgálata történt meg. A modell értékelését a konvergencia és a diszkriminációs validitás szempontjainak értékelésével kezdtem. Az elemzés során vizsgálatra kerültek az indikátorok külső faktorterhelései (*outer loadings*), a konstrukciók belső konzisztenciája, valamint az átlagos magyarázott variancia (*average variance extracted*, AVE) értékei.

Az indikátorok megbízhatóságának és konvergencia validitásának értékelése során az *outer loading* értékeket elemeztem. A szakirodalmi ajánlások szerint a 0,50 feletti

standardizált faktorsúlyok már elfogadható indikátorszintű hozzájárulást jeleznek, míg a magasabb, 0,70 feletti értékek tekinthetők optimálisnak (Hair, 2006).

A konstrukciók belső konzisztenciájának értékelésére a Cronbach-alfa (*Cronbach's alpha*) és a kompozit megbízhatóság (*composite reliability*, CR) mutatók kerültek alkalmazásra. A szakirodalom alapján a 0,70 feletti értékek megfelelő megbízhatóságot jeleznek (Malkanthie, 2015). A konvergens validitás értékelése során az átlagos magyarázott variancia (AVE) került figyelembevételre, amely esetében a Fornell–Larcker-kritérium szerint a 0,50 feletti értékek tekinthetők elfogadhatónak (Fornell–Larcker, 1981).

3. táblázat

A mérési modell megbízhatósági és validitási mutatói
Reliability and validity indicators of the measurement model

Konstrukció	Indikátor	Outer loading	Cronbach-alfa	CR	AVE
Társadalmi befolyás	SI1	0,567	0,847	0,889	0,621
	SI2	0,776			
	SI3	0,864			
	SI4	0,847			
	SI5	0,847			
Hedonikus motiváció	HM1	0,783	0,901	0,931	0,771
	HM2	0,915			
	HM3	0,937			
	HM4	0,869			
Antropomorfizmus	A1	0,857	0,844	0,893	0,678
	A2	0,862			
	A3	0,833			
	A4	0,734			
Bizalom	T1	0,900	0,913	0,933	0,699
	T2	0,825			
	T3	0,890			
	T4	0,715			
	T5	0,816			
	T6	0,855			
Teljesítményelvárás	PE1	0,828	0,814	0,888	0,726
	PE2	0,836			
	PE3	0,890			
Erőfeszítés-elvárás	EE1	0,855	0,815	0,890	0,729
	EE2	0,822			
	EE3	0,883			
Pozitív érzelmek	E1	0,831	0,910	0,937	0,788
	E2	0,882			
	E4	0,904			
	E5	0,929			
Használati hajlandóság	WU1	0,868	0,737	0,845	0,653
	WU2	0,929			
	WU3	0,586			

Forrás: saját szerkesztés.

A mérési modell megbízhatóságának és validitásának értékelése a Cronbach-alfa, a kompozit megbízhatóság, valamint az átlagos magyarázott varianciamutatók alapján történt. Az eredmények alapján minden konstrukció megfelel a megbízhatósági és konvergens validitási követelményeknek. A Cronbach-alfa értékek 0,737 és 0,913 között mozognak, ami a skálák megfelelő belső konzisztenciáját jelzi. A kompozit megbízhatóság minden esetben meghaladja a 0,845-es értéket, így a konstrukciók megbízhatósága magasnak tekinthető. Az AVE értékek szintén minden esetben meghaladják a 0,50-os küszöbértéket (0,621–0,788 tartomány), ami igazolja a konvergens validitás teljesülését.

Az indikátorok külső faktorterhelései (*outer loadings*) általánosságban megfelelő értékeket mutatnak, a legtöbb indikátor 0,70 feletti értéket mutat. Ugyanakkor a használati hajlandóság (*WU*) konstrukció esetében a WU3 indikátor (0,586) és a társadalmi befolyás (*SI*) konstrukció SI1 indikátora (0,567) alacsonyabb faktorterhelést mutatott, azonban elméleti relevanciája miatt a modellben megtartásra került. Összességében a mérési modell megfelel a megbízhatósági és konvergens validitási kritériumoknak, így alkalmas a strukturális modell további elemzésére.

A diszkrimináns validitás vizsgálata során a konstrukciók közötti korrelációk kerültek elemzésre. Az eredményeket a 4. táblázat tartalmazza. A szakirodalmi ajánlások alapján a konstrukciók megfelelően elkülönülnek egymástól, amennyiben a köztük lévő korrelációs értékek nem haladják meg a 0,85-os küszöbértéket (*Henseler et al., 2015*).

Az eredmények alapján a legtöbb konstrukciópár esetében a korrelációs értékek a javasolt határérték alatt maradtak, ami megfelelő diszkrimináns validitást jelez. Azonban egy esetben, a hedonikus motiváció (*HM*) és a használati hajlandóság (*WU*) közötti korreláció enyhén meghaladta a javasolt határértéket (magasabb értéket mutatott [$r = 0,859$]). Ugyanakkor az eredmény a konstrukciók elméleti elkülöníthetősége, valamint a szakirodalmi értelmezési tartomány figyelembevételével elfogadhatónak tekinthető, így a diszkrimináns validitás összességében nem sérült.

4. táblázat

HTMT- (*heterotrait-monotrait*) arányok a diszkrimináns validitás vizsgálatához
HTMT (heterotrait-monotrait) ratios for assessing discriminant validity

	A	T	EE	WU	HM	E	PE	SI
A								
T	0,331							
EE	0,165	0,615						
WU	0,391	0,822	0,768					
HM	0,419	0,634	0,560	0,859				
E	0,384	0,626	0,519	0,710	0,722			
PE	0,397	0,656	0,516	0,788	0,684	0,508		
SI	0,438	0,585	0,573	0,848	0,648	0,432	0,660	

Forrás: saját szerkesztés.

A mérési modell egészét tekintve megállapítható, hogy az alkalmazott konstrukciók megbízhatóan és megfelelő érvényességgel mérik a vizsgált jelenségeket. A konvergens validitási feltételek minden esetben teljesültek, míg a diszkrimináns validitás tekintetében a konstrukciók többsége megfelelő elkülönülést mutat. Ugyanakkor néhány konstrukciópár esetében – különösen a hedonikus motiváció (HM) és a használati hajlandóság (WU) viszonylatában – erős, határértékhez közeli kapcsolat figyelhető meg, ami a konstrukciók közötti szoros elméleti összefüggésre utal, és az eredmények értelmezése során körültekintő interpretációt igényel.

4.2. A strukturális, azaz a külső modell értékelése

A kutatás pilot jellegéből adódóan a strukturális modell elemzésének elsődleges célja a felállított mérési és strukturális modell működésének tesztelése, stabilitásának és értelmezhetőségének vizsgálata volt. A hipotézisek empirikus ellenőrzése ebben a fázisban másodlagos szerepet töltött be, és elsősorban a modell előzetes validálását szolgálta.

A strukturális modell elemzése a látens konstrukciók közötti kapcsolatok vizsgálatán alapult, amely során a standardizált útvonal-együtthatók (*path coefficients*) kerültek értékelésre. A kapcsolatok szignifikanciájának vizsgálatára bootstrap eljárást alkalmaztam, amely az eredeti mintából történő többszöri újra-mintavételezés révén lehetővé tette a paraméterbecslések stabilitásának és megbízhatóságának értékelését.

A strukturális modell értékelése során a következő mutatók kerültek elemzésre: útvonal-koefficiensek (β), t-értékek, p-értékek, konfidenciaintervallumok, valamint a determinációs együtthatók (R^2).

4.2.1. Útvonal-koefficiensek és hipotézisvizsgálat

A bootstrap eljárás eredményei alapján több szignifikáns kapcsolat is azonosítható volt a modellben.

A hedonikus motiváció ($\beta = 0,391$, $p < 0,001$, CI: [0,225; 0,543]) és a társadalmi befolyás ($\beta = 0,302$, $p = 0,001$, CI: [0,134; 0,483]) szignifikáns pozitív hatást gyakorolt a bizalomkonstrukcióra, ami arra utal, hogy az élvezeti tényezők és a társas hatások egyaránt erősítik a bizalom kialakulását.

Az antropomorfizmus és a bizalom között nem volt kimutatható szignifikáns kapcsolat ($\beta = 0,036$, $p = 0,622$, CI: [-0,091; 0,190]), mivel a konfidenciaintervallum a nullát is tartalmazta.

A bizalom jelentős hatást gyakorolt mind az erőfeszítés-elvárásra ($\beta = -0,534$, $p < 0,001$, CI: [-0,696; -0,352]), mind a teljesítményelvárásra ($\beta = 0,582$, $p <$

0,001, CI: [0,449; 0,708]). Az eredmények alapján megállapítható, hogy a bizalom növekedése csökkenti az erőfeszítés észlelt szükségességét, ugyanakkor növeli a rendszer teljesítményével kapcsolatos elvárásokat.

Az erőfeszítés-elvárás negatív hatást gyakorolt a pozitív érzelmekre ($\beta = -0,325$, $p = 0,001$, CI: [-0,527; -0,139]), míg a teljesítményelvárás pozitív kapcsolatot mutatott ugyanezzel a konstrukcióval ($\beta = 0,315$, $p = 0,004$, CI: [0,090; 0,513]).

A pozitív érzelmek kiemelkedően erős hatást gyakoroltak a használati hajlandóságra ($\beta = 0,630$, $p < 0,001$, CI: [0,508; 0,747]), ami a modell egyik legerősebb strukturális kapcsolatának tekinthető.

4.2.2. A modell magyarázóereje (R^2)

A modell magyarázóerejének értékelése az R^2 mutatók alapján történt. Az eredmények alapján a modell mérsékelt magyarázóerővel rendelkezik a vizsgált endogén konstrukciók esetében.

A determinációs együtthatók (R^2) alapján a modell közepes magyarázóerővel rendelkezik. A bizalom ($R^2 = 0,409$), az erőfeszítés-elvárás ($R^2 = 0,285$), a pozitív érzelmek ($R^2 = 0,292$), a teljesítményelvárás ($R^2 = 0,339$), valamint a használati hajlandóság ($R^2 = 0,397$) konstrukciók esetében a modell a variancia jelentős részét képes magyarázni.

4.2.3. Hatásméretetek (f^2)

A hatásméretetek elemzése alapján több erős és közepes hatás is azonosítható volt a modellben. A bizalom erős hatást gyakorolt a teljesítményelvárásra ($f^2 = 0,512$), valamint közepes hatást az erőfeszítés-elvárásra ($f^2 = 0,399$). A pozitív érzelmek és a használati hajlandóság közötti kapcsolat kiemelkedően erős ($f^2 = 0,659$), ami a modell legdominánsabb összefüggését jelzi.

A hedonikus motiváció ($f^2 = 0,156$) és a társadalmi befolyás ($f^2 = 0,092$) a bizalomkonstrukcióra gyakorolt hatása gyengébb, de még értelmezhető mértékű. Az antropomorfizmus hatása ($f^2 = 0,002$) gyakorlatilag elhanyagolható.

Az eredményeket az 5. táblázat foglalja össze.

Összességében megállapítható, hogy a strukturális modell több szignifikáns és elméletileg is értelmezhető kapcsolatot tartalmaz. A legerősebb hatást a pozitív érzelmek és a használati hajlandóság közötti kapcsolat mutatja, míg az antropomorfizmus konstrukció hatása nem bizonyult szignifikánsnak. A modell magyarázóereje közepesnek tekinthető, ami a pilotkutatás jellegével összhangban megfelelő alapot biztosít a további, végleges modellfinomításhoz.

5. táblázat

A strukturális modell hipotézisvizsgálatának eredményei
Results of the structural model hypothesis testing

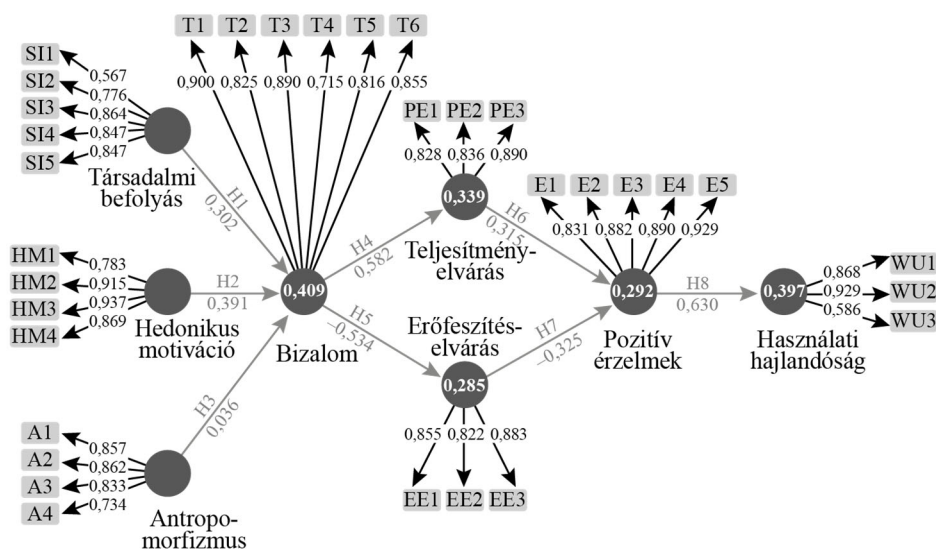
Hipotézis	Kapcsolat	β	t-érték	p-érték	Eredmény
H1	Társadalmi befolyás $\rightarrow$ Bizalom	0,302	3,388	0,001	Elfogadva
H2	Hedonikus motiváció $\rightarrow$ Bizalom	0,391	4,782	< 0,001	Elfogadva
H3	Antropomorfizmus $\rightarrow$ Bizalom	0,036	0,492	0,622	Elutasítva
H4	Bizalom $\rightarrow$ Teljesítményelvárás	0,582	8,712	< 0,001	Elfogadva
H5	Bizalom $\rightarrow$ Erőfeszítés-elvárás	-0,534	6,023	< 0,001	Elfogadva
H6	Teljesítményelvárás $\rightarrow$ Pozitív érzelmek	0,315	2,904	0,004	Elfogadva
H7	Erőfeszítés-elvárás $\rightarrow$ Pozitív érzelmek	-0,325	3,275	0,001	Elfogadva
H8	Pozitív érzelmek $\rightarrow$ Használati szándék	0,630	10,197	< 0,001	Elfogadva

Forrás: saját szerkesztés.

A pilotelemzés eredményei alapján a mérési és strukturális modell stabilnak és értelmezhetőnek bizonyult, így alkalmas a végleges modell további finomítására és nagyobb mintán történő validálására.

2. ábra

Modelleredmények
Model results



Forrás: saját szerkesztés.

5. Diskusszió

A kutatás eredményei alapján megállapítható, hogy a modell több, elméletileg is releváns és empirikusan szignifikáns kapcsolatot azonosított a vizsgált konstrukciók között. A PLS–SEM keretében végzett elemzés igazolta, hogy a bizalom központi szerepet tölt be a felhasználói attitűdök és a viselkedési szándékok kialakulásában.

Az eredmények szerint a hedonikus motiváció és a társadalmi befolyás szignifikáns pozitív hatást gyakorol a bizalomra, ami arra utal, hogy a rendszer használatából fakadó élvezeti tényezők, valamint a társas környezet véleménye egyaránt fontos szerepet játszanak a bizalom kialakulásában. Ez arra enged következtetni, hogy az MI-vel való interakció ne legyen unalmas vagy tisztán funkcionális. A játékos elemek (*gamification*) vagy a szórakoztató, de professzionális stílus növeli a technológiába vetett bizalmat. Azok a felhasználók, akik élvezetesnek találják a rendszert, hajlamosabbak elnézőbbek lenni a technikai nehézségekkel szemben is.

Ezzel szemben az antropomorfizmus hatása nem bizonyult szignifikánsnak, ami arra utalhat, hogy önmagában a rendszer emberihez hasonló tulajdonságainak észlelése nem elegendő a bizalom kialakításához. A menedzsereknek nem érdemes túl sok erőforrást fektetniük abba, hogy az MI-nek „lelket” vagy túlzottan emberi személyiséget adjanak. Az ügyfelek nem várnak el jobb teljesítményt attól, hogy az MI emberibbnek tűnik. Sőt, a túlzott emberszerűség növelheti a kognitív terhelést, mert az ügyfélnek egyszerre kell gépként és emberként is kezelnie a rendszert.

A bizalom kettős hatása jól kirajzolódik az elvárások rendszerében: míg a teljesítményelvárás növekedéséhez vezet, addig az erőfeszítés-elvárás csökkenti. Ez arra utal, hogy a magasabb bizalom a rendszer hatékonyságának erősebb észlelésével jár együtt, ugyanakkor csökkenti a használathoz kapcsolódó kognitív vagy gyakorlati nehézségek észlelését. Ha az ügyfél bízik a rendszerben, kevésbé fogja azt bonyolultnak vagy nehézkesnek érezni. A bizalom építéséhez elengedhetetlen az átláthatóság: világossá kell tenni, mikor kommunikál az ügyfél MI-vel, és mit kezd a rendszer az adataival. A szolgáltatóknak hangsúlyozniuk kell a rendszer hibátlanágát és következetességét az emberi ügyintézőkkel szemben, mert ez közvetlenül javítja a teljesítményelvárást.

Az elvárások hatása a felhasználói érzelmekre szintén koherens mintázatot mutat. A teljesítményelvárás pozitívan, míg az erőfeszítés-elvárás negatívan befolyásolja a pozitív érzelmeket, ami összhangban van az elvárásalapú technológiaelfogadási modellekkel. Ez azt jelenti, hogy a felhasználók érzelmi reakciói jelentős mértékben függenek attól, hogy a rendszert mennyire hatékynak, illetve mennyire erőfeszítés-igényesnek érzik.

A modell legerősebb kapcsolatát a pozitív érzelmek és a használati szándék között azonosítottam, ami azt jelzi, hogy az érzelmi dimenzió kulcsszerepet játszik a technológia elfogadásában. Ez az eredmény megerősíti, hogy a felhasználói szándékok nem kizárólag racionális értékelésen, hanem erőteljes affektív tényező-kön is alapulnak. Azaz a kampányoknak nem csak a technikai paraméterekre (pl. 0–24 elérhetőség) kellene fókuszálniuk, hanem arra, hogy az MI használata nyugalmat, elégedettséget és gördülékenységet biztosítson az ügyfeleknek. A kezelő-felületek kialakításánál pedig elsődleges szempont legyen a „kellemes élmény” (színek, hangnem, válaszsebesség), mivel a használati szándék az eredmények alapján inkább érzelmi, mint tisztán racionális alapú.

Összességében a pilot vizsgálat eredményei azt mutatják, hogy a modell jól működik, és a konstrukciók közötti kapcsolatok elméletileg értelmezhetők és empirikusan többnyire alátámasztott struktúrát alkotnak.

Irodalom

- Balmer, R. E. – Levin, S. L. – Schmidt, S. (2020): Artificial intelligence applications in telecommunications and other network industries. *Telecommunications Policy*, 44(6), 6.
<https://doi.org/10.1016/j.telpol.2020.101977>
- Chi, O. H. – Chi, C. G. – Gursoy, D. – Nunkoo, R. (2023): Customers' acceptance of artificial intelligent service robots: the influence of trust and culture. *International Journal of Information Management*, 55, 416–424.
- Chow, C. S. K. – Zhan, G. – Wang, H. – He, M. (2023): Artificial intelligence (AI) adoption: an extended compensatory level of acceptance. *Journal of Electronic Commerce Research*, 24(1), 84–106.
- Corkrey, R. – Parkinson, L. (2002): Interactive voice response: review of studies 1989–2000. *Behavior Research Methods Instruments & Computers*, 34(3), 342–353.
<https://doi.org/10.3758/BF03195462>
- Davis, F. D. (1986): *A technology acceptance model for empirically testing New End-user information systems: theory and result*. Doctoral dissertation, MIT Sloan School of Management, Cambridge, MA.
- Della Corte, V. – Sepe, F. – Gursoy, D. – Prisco, A. (2023): Role of trust in customer attitude and behaviour formation towards social service robots. *International Journal of Hospitality Management*, 114(2). <https://doi.org/10.1016/j.ijhm.2023.103587>
- Duarte, V. – Zuniga-Jara, S. – Contreras, S. (2022): Machine learning and marketing: a systematic literature review. *Institute of Electrical and Electronics Engineers*, 10, 93273–93288.
<https://doi.org/10.1109/ACCESS.2022.3202896>
- Fornell, C. – Larcker, D. F. (1981): Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39.
<https://doi.org/10.2307/3151312>
- Gursoy, D. – Chi, H. O. – Lu, L. – Nunkoo, R. (2019): Consumers acceptance of artificially intelligent (AI) device use in service delivery. *International Journal of Information Management*, 49, 157–169. <https://doi.org/10.1016/j.ijinfomgt.2019.03.008>

- Haenlein, M. – Kaplan, A. M. (2004): A beginner's guide to partial least squares analysis. *Understanding Statistics*, 3(4), 283–297. https://doi.org/10.1207/s15328031us0304_4
- Hair, J. F. (ed.) (2006): *Multivariate data analysis* (6th ed). Pearson Prentice Hall.
- Kazár, K. (2014): A PLS-útelemzés és alkalmazása egy márkaközösség pszichológiai érzetének vizsgálatára. *Statistikai Szemle*, 91(1), 33–52.
- Keszei, T. – Zsukk, J. (2017): Az új technológiák fogyasztói elfogadása. A magyar és nemzetközi szakirodalom áttekintése és kritikai értékelése. *Vezetéstudomány – Budapest Management Review*, 48(10), 38–47. <https://doi.org/10.14267/VEZTUD.2017.10.05>
- Khan, A. N. – Jabeen, F. – Mehmood, K. – Ali Soomro, M. – Bresciani, S. (2023): Paving the way for technological innovation through adoption of artificial intelligence in conservative industries. *Journal of Business Research*, 165(10), 3–10. <https://doi.org/10.1016/j.jbusres.2023.114019>
- Kim, M. K. – Park, M. C. – Jeong, D. H. (2004): The effect of customer satisfaction and switching barrier on customer loyalty in Korean mobile telecommunication services. *Telecommunications Policy*, 28(2), 145–159. <https://doi.org/10.1016/j.telpol.2003.12.003>
- Lacity, M. C. – Willcocks, L. P. (2016): A new approach to automating services. *MIT Sloan Management Review*, 58(1), 41–49.
- Lee, S. M. – Lee, D. (2020): „Untact”: a new customer service strategy in the digital age. *Service Business*, 14(1), 1–22. <https://doi.org/10.1007/s11628-019-00408-2>
- Li, W. – Ding, H. – Gui, J. – Tank, Q. (2024): Patient acceptance of medical service robots in the medical intelligence era: an empirical study based on an extended AI device use acceptance model. *Humanities and Social Sciences Communications*, 11(1), 1–13. <https://doi.org/10.1057/s41599-024-04028-8>
- Lin, H. – Chi, O. H. – Gursoy, D. (2020): Antecedents of customers' acceptance of artificially intelligent robotic device use in hospitality services. *Journal of Hospitality Marketing and Management*, 29(5), 530–549. <https://doi.org/10.1080/19368623.2020.1685053>
- Lu, V. N. – Wirtz, J. – Kunz, H. W. – Paluch, S. – Gruber, T. – Martins, A. – Patterson, G. P. (2020): Service robots, customers and service employees: what can we learn from the academic literature and where are the gaps? *Journal of Service Theory and Practice*, 30(3), 361–391. <https://doi.org/10.1108/jstp-04-2019-0088>
- Malkanthe, A. (2015): *Structural equation modeling with AMOS*. <https://doi.org/10.13140/RG.2.1.1960.4647>
- Manzoor, S. R. – Ullah, R. – Khattak, A. – Ullah, M. – Han, H. (2024): Exploring tourist perceptions of artificial intelligence devices in the hotel industry: impact of industry 4.0. *Journal of Travel and Tourism Marketing*, 41(2), 272–291. <https://doi.org/10.1080/10548408.2024.2310169>
- Mittelstadt, B. (2019): Principles alone cannot guarantee ethical AI. *Nature Machine Intelligence*, 1, 501–507. <https://doi.org/10.1038/s42256-019-0114-4>
- Necz, D. (2023): A mesterséges intelligencia felhasználásával történő adatkezelések egyes sajátos szempontjai. *Acta Humana – Emberi Jogi Közlemények*, 10(3), 95–123. <https://doi.org/10.32566/ah.2022.3.4>
- Nemzeti Média- és Hírközlési Hatóság (NMHH) (2025): *Mobilpiaci jelentés – 2025. II. negyedév*. Budapest: NMHH. [Mobilpiaci jelentés – 2025. II. félév – Nemzeti Média- és Hírközlési Hatóság](https://www.nmhh.hu/mobilpiaci-jeleltes-2025-ii-felév)
- Ouyang, Y. – Wang, L. – Yang, A. – Shah, M. – Belanger, D. – Gao, T. – Wei, L. – Zhang, Y. (2021): The next decade of telecommunications artificial intelligence. *Telecommunications Science*, 37(3), 1–36. <https://doi.org/10.26599/air.2022.9150003>

- Peng, C. – van Doorn, J. – Eggers, F. – Wieringa, J. E. (2022): The effect of required warmth on consumer acceptance of artificial intelligence in service: the moderating role of AI-human collaboration. *International Journal of Information Management*, 66, 1–11.
<https://doi.org/10.1016/j.ijinfomgt.2022.102533>
- Pininti, S. (2024): Scalable AI and ML-powered customer support for Telecom enterprises: optimizing engagement for 200 million + customers. *International Journal of Intelligent Systems and Applications in Engineering*, 12(23s), 3616–3617.
<https://ijisae.org/index.php/IJISAE/article/view/7800>
- Ribeiro, M. A. – Gursay, D. – Chi, O. H. (2022): Customer acceptance of autonomous vehicles in travel and tourism. *Journal of Travel Research*, 61(3), 620–636.
<https://doi.org/10.1177/0047287521993578>
- Roy, P. – Ramaprasad, B. S. – Chakraborty, M. – Prabhu, N. – Rao, S. (2024): Customer acceptance of use of artificial intelligence in hospitality service: an Indian hospitality sector perspective. *Global Business Review*, 25(3), 832–851. <https://doi.org/10.1177/0972150920939753>
- Santosh, E. – Rao, U. V. A. – Sravan, E. K. (2020): Application of artificial intelligence in improving operational efficiency in telecom industry. *International Journal on Emerging Technologies*, 11(3), 65–69.
- Schiavo, G. – Businaro, S. – Zancanaro, M. (2024): Comprehension, apprehension, and acceptance: understanding the influence of literacy and anxiety on acceptance of artificial intelligence. *Technology in Society*, 77(3). <https://doi.org/10.1016/j.techsoc.2024.102537>
- Sen, R. – Vasudevan, S. – Britto, R. – Prasath, M. (2024): Ascertaining trustworthiness of AI systems in telecommunications. *ITU Journal on Future and Evolving Technologies*, 5(4), 503–514.
<https://doi.org/10.52953/WIBX7049>
- Sharma, A. (2021): Artificial intelligence in healthcare. *International Journal of Research in Humanities, Arts, Medicine and Science*, 5(1), 106–109.
- Tisóczki, J. (2022): A mesterséges intelligencia alkalmazása az egészségügyi ellátási folyamatokban. *Biztonságtudományi Szemle*, 4(2), 137–153.
- Umamaheswari, S. – Valarmathi, A. (2023): Role of artificial intelligence in the banking sector. *Journal of Survey in Fisheries Sciences*, 10(4S), 2841–2849.
- Venkatesh, V. – Morris, M. G. – Davis, G. B. – Davis, F. D. (2003): User acceptance of information technology: toward a unified view. *MIS Quarterly*, 27(3), 425–478.
<https://doi.org/10.2307/30036540>
- Wang, L. – Huang, N. – Hong, Y. – Liu, L. – Guo, X. – Chen, G. (2023): Voice-based AI in call center customer service: a natural field experiment. *Production and Operations Management*, 32(1), 1–17. <https://doi.org/10.1111/poms.13953>
- Wirtz, J. – Patterson, P. G. – Kunz, W. H. – Gruber, T. – Lu, V. N. – Paluch, S. – Martins, A. (2018): Brave new world: service robots in the frontline. *Journal of Service Management*, 29(5), 907–931. <https://doi.org/10.1108/JOSM-04-2018-0119>
- Wong, I. A. – Zhang, T. – Lin, Z. C. – Peng, Q. (2023): Hotel AI service: are employees still needed? *Journal of Hospitality and Tourism Management*, 55, 416–424.
<https://doi.org/10.1016/j.jhtm.2023.05.005>