

Sveučilište Josipa Jurja Strossmayera u Osijeku

Ekonomski fakultet u Osijeku

Doktorski studij Management

Bruno Budimir

**ODREDNICE PRIHVAĆANJA I KORIŠTENJA
GENERATIVNE UMJETNE INTELIGENCIJE MJERENE
PROŠIRENIM UTAUT2 MODELOM**

Doktorski rad

Osijek, 2025.

Sveučilište Josipa Jurja Strossmayera u Osijeku
Ekonomski fakultet u Osijeku
Doktorski studij Management

Bruno Budimir

**ODREDNICE PRIHVAĆANJA I KORIŠTENJA
GENERATIVNE UMJETNE INTELIGENCIJE MJERENE
PROŠIRENIM UTAUT2 MODELOM**

Doktorski rad

Matični broj studenta: 402

E-mail: budimir@efos.hr

Mentor: prof. dr. sc. Antun Biloš

Osijek, 2025.

Josip Juraj Strossmayer University of Osijek
Faculty of Economics and Business in Osijek
Doctoral study Management

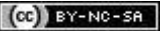
Bruno Budimir

**DETERMINANTS OF ACCEPTANCE AND USE OF
GENERATIVE ARTIFICIAL INTELLIGENCE MEASURED
BY THE EXTENDED UTAUT2 MODEL**

Doctoral thesis

Osijek, 2025.

IZJAVA
O AKADEMSKOJ ČESTITOSTI, PRAVU PRIJENOSA
INTELEKTUALNOG VLASNIŠTVA, SUGLASNOSTI ZA OBJAVU
U INSTITUCIJSKIM REPOZITORIJIMA I ISTOVJETNOSTI
DIGITALNE I TISKANE VERZIJE RADA

1. Kojom izjavljujem i svojim potpisom potvrđujem da je doktorski rad isključivo rezultat osobnoga rada koji se temelji na mojim istraživanjima i oslanja se na objavljenu literaturu. Potvrđujem poštivanje nepovredivosti autorstva te točno citiranje radova drugih autora i referiranje na njih.
2. Kojom izjavljujem da je Ekonomski fakultet u Osijeku, bez naknade u vremenski i teritorijalno neograničenom opsegu, nositelj svih prava intelektualnoga vlasništva u odnosu na navedeni rad pod licencom *Creative Commons Imenovanje – Nekomercijalno – Dijeli pod istim uvjetima 3,0 Hrvatska*. 
3. Kojom izjavljujem da sam suglasan/suglasna da se trajno pohrani i objavi moj rad u institucijskom digitalnom repozitoriju Ekonomskoga fakulteta u Osijeku, repozitoriju Sveučilišta Josipa Jurja Strossmayera u Osijeku te javno dostupnom repozitoriju Nacionalne i sveučilišne knjižnice u Zagrebu (u skladu s odredbama Zakona o visokom obrazovanju i znanstvenoj djelatnosti, (NN 119/2022).
4. Izjavljujem da sam autor/autorica predanog rada i da je sadržaj predane elektroničke datoteke u potpunosti istovjetan sa dovršenom tiskanom verzijom rada predanom u svrhu obrane istog.

Ime i prezime studenta/studentice: *Bruno Budimir*

Matični broj studenta: 402

OIB: 44188346757

e-mail za kontakt: *budimir@efos.hr*

Naziv studija: *Doktorski studij Management*

Naslov rada: *Odrednice prihvaćanja i korištenja generativne umjetne inteligencije mjerene proširenim UTAUT2 modelom*

Mentor rada: *prof. dr. sc. Antun Biloš*

U Osijeku, lipanj 2025. godine



Bruno Budimir

TEMELJNA DOKUMENTACIJSKA KARTICA

Sveučilište Josipa Jurja Strossmayera u Osijeku
Ekonomski fakultet u Osijeku

Doktorski rad

Znanstveno područje: Društvene znanosti

Znanstveno polje: Ekonomija

ODREDNICE PRIHVAĆANJA I KORIŠTENJA GENERATIVNE UMJETNE INTELIGENCIJE MJERENE PROŠIRENIM UTAUT2 MODELOM

Doktorski rad izrađen je u: Osijeku

Mentor: prof. dr. sc. Antun Biloš

Kratki sadržaj doktorske disertacije:

Doktorski rad istražuje odrednice prihvaćanja i korištenja generativne umjetne inteligencije kroz mjerene proširenim UTAUT2 modelom. Navedeni teorijski okvir proširen je konstruktima osobne inovativnosti i povjerenja, a čestice konstrukta korištenja modificirane i prilagođene kontekstu istraživanja. Istraživanje je provedeno u Ujedinjenom Kraljevstvu i Republici Hrvatskoj na kvotnom uzorku od 400 ispitanika po tržištu. Rad donosi i teorijske preporuke za buduća istraživanja na ovu i srodne teme te navodi važne doprinose za praktičare u čijem je interesu razumijevanje korisnika ove napredne tehnologije.

Broj stranica: 259

Broj slika: 23

Broj tablica: 41

Broj literaturnih navoda: 343

Jezik izvornika: Hrvatski

Ključne riječi: UTAUT2 model, Generativna umjetna inteligencija, Umjetna inteligencija, Prihvaćanje tehnologije, Proširenje modela

Datum obrane:

Stručno povjerenstvo za obranu:

Doktorski rad pohranjen je u: Nacionalnoj i sveučilišnoj knjižnici Zagreb, Ul. Hrvatske bratske zajednice 4, Zagreb; Gradskoj i sveučilišnoj knjižnici Osijek, Europska avenija 24, Osijek; Sveučilištu Josipa Jurja Strossmayera u Osijeku, Trg. sv. Trojstva 3, Osijek; Ekonomskom fakultetu u Osijeku, Trg Lj. Gaja 7, Osijek

BASIC DOCUMENTATION CARD

Josip Juraj Strossmayer University of Osijek
Faculty of Economics and Business in Osijek

PhD thesis

Scientific Area: Social sciences

Scientific Field: Economics

DETERMINANTS OF ACCEPTANCE AND USE OF GENERATIVE ARTIFICIAL INTELLIGENCE MEASURED BY THE EXTENDED UTAUT2 MODEL

Thesis performed at: Osijek

Supervisor: Antun Biloš, PhD, Full Professor

Short abstract:

The doctoral thesis explores the determinants of acceptance and use of generative artificial intelligence measured through the extended UTAUT2 model. The theoretical framework was expanded with the constructs of personal innovativeness and trust, while the items related to USE construct were modified and adapted to the research context. The study was conducted in the United Kingdom and the Republic of Croatia on a quota sample of 400 respondents per market. The measurement model was tested using PLS-SEM analysis. The dissertation also offers theoretical recommendations for future research on this and related topics and provides practical implications for professionals interested in understanding users of this advanced technology.

Number of pages: 259

Number of figures: 23

Number of tables: 41

Number of references: 343

Original in: Croatian

Keywords: UTAUT2, Generative artificial intelligence, Artificial intelligence (AI), Technology Acceptance, Model Extension

Date of the thesis defense:

Reviewers:

Thesis deposited in: National and University Library in Zagreb, Ul. Hrvatske bratske zajednice 4, Zagreb; City and University Library of Osijek, Europska avenija 24, Osijek; Josip Juraj Strossmayer University of Osijek, Trg. sv. Trojstva 3, Osijek; Faculty of Economics and Business in Osijek, Trg Lj. Gaja 7, Osijek

SAŽETAK

Pojava generativne umjetne inteligencije (Gen AI) dovela je do naglog porasta interesa opće javnosti za ovom tehnologijom, posebno od kraja studenog 2022. godine kada je ChatGPT, kao prvi takav napredni alat, lansiran u javnost. Nagli rast i razvoj te sveprisutna upotreba takvih alata stvorili su značajnu potrebu za dubljim razumijevanjem čimbenika koji utječu na njihovo usvajanje od strane korisnika, kako kod istraživača u akademskoj zajednici tako i u široj javnosti. Premda se teorijski modeli prihvaćanja tehnologije koriste već desetljećima, njihova primjena u kontekstu generativne umjetne inteligencije relativno je ograničena. Upravo iz tog razloga ovaj doktorski rad koristi jedan od najmodernijih i najkorištenijih modela prihvaćanja tehnologije - Ujedinjenu teoriju o prihvaćanju i korištenju tehnologije 2 (UTAUT2 - engl. *Unified Theory of Acceptance and Use of Technology 2*). Navedeni teorijski okvir za istraživanja prihvaćanje tehnologije vrlo je fleksibilan za modifikacije u svrhu prilagođavanja modela za određeni kontekst istraživanja. Shodno tome, svrha ovog dokorskog rada je analizirati i teorijski proširiti UTAUT2 model za istraživanja u kontekst generativne umjetne inteligencije. Proširenje UTAUT2 modela u ovom radu uključuju uvođenje dva nova konstrukta – osobnu inovativnost i povjerenje, ali se u svrhu boljeg razumijevanja korištenja mijenjaju i čestice konstrukta stvarnog korištenja. Predloženi konceptualni model, u konačnici, temelji se na konstruktima očekivanog učinka (PE), očekivanih napora (EE), društvenog utjecaja (SI), olakšavajućih uvjeta (FC), hedonističke motivacije (HM), vrijednosti cijene (PV), navike (HT), osobne inovativnosti (PI) i povjerenja (TR) za koje se pretpostavlja da su prediktori namjere ponašanja (BI) i stvarnog korištenja (USE). Osim navedenih nezavisnih i zavisnih varijabli, rad mjeri i utjecaj dobi, spola i iskustva kao moderatora navedenih odnosa.

U empirijskom dijelu rada istraživanje je provedeno na dva tržišta, u Ujedinjenom Kraljevstvu kao naprednoj ekonomiji po IMF-ovom Indeksu spremnosti na umjetnu inteligenciju i u Republici Hrvatskoj kao ekonomiji u razvoju po navedenom indeksu. Ukupni uzorak uključuje 802 ispitanika starijih od 18 godina, koristeći kvotni uzorak po dobi i spolu, ujednačen za oba tržišta s ciljem bolje usporedne analize dvaju tržišta. Istraživanje koristi visokostrukturirani istraživački instrument temeljen na postojećim mjernim ljestvicama, a za testiranje postavljenih hipoteza i analizu utjecajnih čimbenika korišteno je modeliranje strukturnih jednadžbi metodom parcijalnih najmanjih kvadrata (PLS-SEM). Znanstveni doprinos ovog rada očituje se u konceptualnom, metodološkom i aplikativnom smislu. Konceptualno, rad doprinosi boljem razumijevanju teorijskog okvira UTAUT2 modela i prilagođava ga za istraživanja u kontekstu generativne umjetne inteligencije. Metodološki gledano, doprinos ovog dokorskog rada

rezultati su testiranja tog prilagođenog modela, ali i uključivanje šireg demografskog spektra i moderatora koji su u ovim istraživanjima prečesto izostavljeni. U konačnici, aplikativni doprinos leži u pružanju konkretnih uvida marketinškim stručnjacima, tvorcima politika, pojedincima i poslovnim subjektima za poticanje šire i uspješnije integracije generativne umjetne inteligencije u različite scenarije uporabe.

Ključne riječi: generativna umjetna inteligencija; prihvatanje tehnologije; proširenje modela; umjetna inteligencija; UTAUT2

SUMMARY

Public interest in generative artificial intelligence (Gen AI) has grown quickly as it became widely known, particularly since late November 2022, when ChatGPT, the first such advanced tool, was launched for public use. The swift growth, development and widespread application of these tools have created a significant need for a deeper understanding of the determinants that influence their adoption by users, both within academia and the general public. Although theoretical models of technology acceptance have been used for decades, their application in the context of generative artificial intelligence remains relatively limited. For this reason, this doctoral dissertation employs one of the most modern and widely used models of technology acceptance – Unified Theory of Acceptance and Use of Technology 2, also known as UTAUT2. This theoretical framework is highly flexible and can be adapted to various research contexts. Accordingly, the purpose of this PhD thesis is to analyze and theoretically extend the UTAUT2 model for research in the context of generative artificial intelligence. The extended UTAUT2 model in this research incorporates two new constructs – personal innovativeness and trust – and modifies the measurement items for actual usage to better capture user behavior. The proposed conceptual model is based on the constructs of performance expectancy (PE), effort expectancy (EE), social influence (SI), facilitating conditions (FC), hedonic motivation (HM), price value (PV), habit (HT), personal innovativeness (PI), and trust (TR), which are hypothesized to be predictors of behavioral intention (BI) and actual usage (USE). In addition to the independent and dependent variables, the model also examines the moderating effects of age, gender, and experience of users.

In the empirical part of the thesis, the research was conducted across two markets: the United Kingdom, representing a developed economy according to the IMF's Artificial Intelligence Readiness Index, and the Republic of Croatia, representing an emerging economy. The total sample includes 802 respondents over the age of 18, using a quota sampling method based on age and gender, balanced across both markets to allow for a better comparative analysis of the two markets. The research uses a highly structured questionnaire based on validated measurement scales. For hypothesis testing and analysis of influencing factors, structural equation modeling using the partial least squares method (PLS-SEM) was employed. The scientific contribution of this dissertation is threefold: conceptual, methodological, and applicative. Conceptually, the research advances the understanding of the UTAUT2 model and adapts it to the Gen AI context. Methodologically, it contributes by testing the extended model across a broader demographic spectrum and including moderators that are often dropped in

similar studies. Finally, the applicative contribution lies in providing actionable insights for marketing professionals, policymakers, individuals, and organizations to support the broader and more effective integration of generative artificial intelligence across various usage scenarios.

Keywords: UTAUT2, Generative Artificial Intelligence, Artificial Intelligence (AI), Technology Acceptance, Model Extension

SADRŽAJ

SAŽETAK.....	1
SUMMARY.....	3
1. UVOD.....	8
1.1. Problem istraživanja	10
1.2. Svrha i ciljevi istraživanja	11
1.3. Hipoteze istraživanja	12
1.4. Metodologija istraživanja	14
1.5. Očekivani znanstveni doprinos.....	16
1.6. Struktura rada	17
2. UMJETNA INTELIGENCIJA.....	19
2.1. Razvoj umjetne inteligencije.....	21
2.2. Definiranje umjetne inteligencije	29
2.3. Generativna umjetna inteligencija	33
2.4. Primjena generativne umjetne inteligencije.....	38
2.5. Etička perspektiva uporabe generativne umjetne inteligencije	44
3. MODELI PRIHVATANJA TEHNOLOGIJE	46
3.1. UTAUT	46
3.1.1. Razvoj modela.....	46
3.1.2. Konstrukti modela UTAUT	53
3.1.2.1. PE – Očekivani učinak.....	54
3.1.2.2. EE – Očekivani naponi.....	57
3.1.2.3. SI – Društveni utjecaj.....	59
3.1.2.4. FC – Olakšavajući uvjeti	62
3.2. UTAUT2	65
3.2.1. Konstrukti modela UTAUT2	65
3.2.1.1. HM – Hedonistička motivacija	66
3.2.1.2. PV – Vrijednost cijene	67
3.2.1.3. HT – Navika	68
3.3. Promjene unutar UTAUT2 modela u odnosu na izvorni UTAUT.....	69
4. PREGLED LITERATURE.....	74
4.1. Modifikacije i proširenja UTAUT2 modela	91
4.1.1. Stvarno korištenje (USE)	91

4.1.2.	Osobna inovativnost – (PI <i>Personal innovativeness</i>).....	97
4.1.3.	Sigurnost i rizici	98
4.1.4.	Povjerenje (TR – <i>Trust</i>).....	100
4.2.	Moderatori	103
4.3.	Analiza prethodnih istraživanja.....	106
4.3.1.	Odnosi u prethodnim istraživanjima	109
4.3.2.	Preliminarno istraživanje.....	118
5.	USPOREDBA REPUBLIKE HRVATSKE I UJEDINJENOG KRALJEVSTVA PO SPREMNOSTI NA UMJETNU INTELIGENCIJU	122
5.1.	Upotreba generativne umjetne inteligencije u Hrvatskoj.....	125
5.2.	Upotreba generativne umjetne inteligencije u Ujedinjenom Kraljevstvu.....	126
6.	KONCEPTUALNI DIZAJN ISTRAŽIVANJA	128
6.1.	Instrument istraživanja.....	129
6.2.	Prošireni UTAUT2 model	134
6.3.	Hipoteze	135
7.	EMPIRIJSKO ISTRAŽIVANJE	137
7.1.	Metodologija primarnog istraživanja	137
7.2.	Proces prikupljanja podataka	141
7.2.1.	Ukupni uzorak	143
7.3.	Rezultati primarnog istraživanja	145
7.3.1.	Testiranje hipoteza	153
7.4.	Analiza rezultata s poduzorka iz Ujedinjenog Kraljevstva	161
7.4.1.	Prikupljanje podataka i uzorak na tržištu Ujedinjenog Kraljevstva.....	161
7.4.2.	Rezultati primarnog istraživanja na tržištu Ujedinjenog kraljevstva	163
7.4.2.1.	Testiranje odnosa u modelu na uzorku iz Ujedinjenog Kraljevstva	169
7.5.	Analiza rezultata s poduzorka iz Republike Hrvatske.....	175
7.5.1.	Prikupljanje podataka i uzorak na tržištu Republike Hrvatske	175
7.5.2.	Rezultati primarnog istraživanja na tržištu Republike Hrvatske.....	177
7.5.2.1.	Testiranje odnosa u modelu na uzorku iz Republike Hrvatske	183
8.	RASPRAVA.....	189
8.1.	Metodološka razmatranja i evaluacija mjernog modela	189
8.2.	Sinteza glavnih nalaza i usporedba s prethodnim istraživanjima	193
8.3.	Implikacije istraživanja	201
8.3.1.	Teorijske implikacije istraživanja	201

8.3.2.	Praktične implikacije istraživanja	203
8.4.	Usporedba konceptualnog modela s prethodnicima	205
8.5.	Ograničenja rada	207
8.6.	Preporuke za buduća istraživanja	209
8.7.	Osvrt na korištenje generativne umjetne inteligencije u pisanju doktorskog rada... ..	213
9.	ZAKLJUČAK	215
10.	POPIS LITERATURE	219
11.	POPIS TABLICA.....	254
12.	POPIS SLIKA	255
13.	POPIS KRATICA.....	256
14.	POPIS OBJAVLJENIH RADOVA PRISTUPNIKA	258

1. UVOD

Premda pojam umjetne inteligencije nije novitet među stručnjacima i akademskim istraživačima koji se bave prihvatanjem i korištenjem tehnologije, pojavom alata generativne umjetne inteligencije navedeno je područje doživjelo snažan rast interesa kod akademskih istraživača, ali i javnosti za upotrebu navedenih alata. Kako navodi Bloomberg (2023), tržište generativne umjetne inteligencije do 2032. godine vrijedit će otprilike 1,3 bilijuna američkih dolara.

Generativna umjetna inteligencija relativno je novo polje unutar umjetne inteligencije, a koja je snažno privukla pozornost javnosti u studenom 2022. godine kada je tvrtka OpenAI javnosti predstavila ChatGPT (Brühl, 2023). Generativna umjetna inteligencija (engl. *generative artificial intelligence* - Gen AI) koristi generativno modeliranje i napredak u dubokom učenju (engl. *deep learning*) kako bi generirala raznovrstan sadržaj u velikom obujmu koristeći postojeće medije poput teksta, zvuka, videa, slika i drugih grafika (Jovanović & Campbell, 2022).

Za razliku od klasičnih modela strojnog učenja unutar umjetne inteligencije, koji prepoznaju uzorke u podacima za obuku kako bi na temelju njih naučili predviđati, klasificirati, davati personalizirane preporuke ili pružati podršku u donošenju odluka, generativna umjetna inteligencija omogućuje izrazito brzo stvaranje novih sadržaja na temelju korisničkog upita (Murugesan i Cherukuri, 2023). Generativna umjetna inteligencija ima sposobnost učenja iz postojećih podataka na temelju kojih generira nove sadržaje. Rezultat generiranja često je vrlo realističan sadržaj koji zadržava karakteristike izvornog skupa podataka, ali ne ponavlja sadržaj, odnosno daje uvijek jedinstven rezultat (Gartner, 2023; Rois-Campos i sur., 2023).

Rad pod nazivom „*Consumer acceptance and use of information technology: Extending the unified theory of acceptance and use of technology*“ (hrv. Potrošačko prihvatanje i korištenje informacijske tehnologije: proširenje ujedinjene teorije prihvatanja i korištenja tehnologije) u kojem su Venkatesh i sur. (2012) predstavili UTAUT2 model danas je jedan od najvažnijih teorijskih okvira za praktičare i akademske istraživače čiji je cilj ispitivanje prihvatanja i korištenja tehnologija. Navedeni rad važan je u smislu teorijske perspektive za razumijevanje pitanja povezanih s prihvatanjem tehnologije u različitim kontekstima, bilo samostalno ili u kombinaciji s drugim teorijama i dodatnim vanjskim varijablama (Tamilmani i sur., 2021).

Venkatesh i sur. (2012) navode kako je u izvornom radu u kojem je predstavljen UTAUT2, a u kojem se ispituje prihvatanje i korištenje mobilnog interneta, većina hipoteza podržana.

UTAUT2 sa samo izravnim učincima objašnjava 44 % varijance u namjeri ponašanja dok UTAUT2 koji uključuje interakcijske odnose objašnjava 74 % varijance u namjeri ponašanja. Kada se promatra objašnjenje korištenja tehnologije, model UTAUT sa samo izravnim učincima objašnjava 35 % varijance dok s uključenim interakcijskim pojmovima objašnjava 52 % varijance. Sve to predstavlja znatno unaprjeđenje u odnosu na izvorni UTAUT.

U ranijim studijama koje su radile detaljne preglede razvoja ovog modela, uočeno je kako 77 % znanstvenih i stručnih članaka citira UTAUT2 u opće svrhe, a preostalih 23 % to čini u kombinaciji s vanjskim teorijama i rijetkim uključivanjem moderatora (Tamilmani i sur., 2017; Tamilmani i sur., 2021). Tamilmani i sur. (2021) navode kako su istraživači u svojim nastojanjima da razumiju individualno prihvaćanje i korištenje tehnologije primijenili, integrirali i proširili UTAUT2 u raznim kontekstima, a što se može podijeliti u šest kategorija:

- različite vrste korisnika
- različite vrste organizacija
- različite vrste tehnologija
- različite vrste zadataka
- različita vremena
- različite lokacije.

Marikyan i Papagiannidis (2023) navode kako je UTAUT tijekom godina pokazao široku primjenu čime je povećana mogućnost generaliziranja teorije jer su tijekom godina znanstvenici proširivali model kako bi ga prilagodili kontekstu i poboljšali njegovu prediktivnu snagu. UTAUT2 izvorno nije ni dizajniran s posebnim naglaskom na određenu novu tehnologiju ili geografsku lokaciju. Cilj teorije bio je predstaviti sveobuhvatan okvir za ispitivanje prihvaćanja i korištenja tehnologije, a navedena proširenja u konceptualnom dizajnu modela koriste se za povećanje preciznosti u objašnjavanju ponašanja korisnika ovisno o kontekstu onoga što se istražuje (Marikyan i Papagiannidis, 2023; Venkatesh i sur., 2012).

Ujedinjena teorija prihvaćanja i korištenja tehnologije jedna je od najcitiranijih teorija koja služi kao model za istraživanja iz područja prihvaćanja i korištenja informacijskih tehnologija. Izvorni rad u kojem su Venkatesh i sur. (2003) predstavili UTAUT, u studenom 2024. godine, prema Google Znalcu, ima više od 56 tisuća citata. S druge strane, UTAUT2 koji su predstavili Venkatesh i sur. (2012) ima više od 18 tisuća citata. Ukoliko se tome pridodaju sestrinski modeli koji su prethodnici UTAUT2 modela, kao što su TAM2 (Venkatesh i Davis, 2000) koji ima više od 33 tisuće citata ili TAM3 (Venkatesh i sur., 2008) koji ima više od 10 tisuća citata,

dolazi se do ukupne brojke od 117 tisuća citata samo kod modela za ispitivanje prihvatanja i korištenja tehnologija, a u čijim je kreacijama sudjelovao profesor Viswanath Venkatesh, najcitiraniji svjetski znanstvenik u ovoj domeni (Venkatesh, 2024).

Uzevši u obzir navedeno, autor ovoga rada nije u mogućnosti za pripremu ovog dokumenta uključiti sva dostupna istraživanja jer im se broj povećava gotovo svakodnevno, ali je bilo nužno analizirati sve relevantne radove do određenog trenutka. U pregledu literature, vidljivog u Tablici 1, navedeno je 68 istraživanja relevantnih za temu predloženog doktorskog rada, a koji su objavljeni do srpnja 2024. godine.

1.1. Problem istraživanja

U posljednjih nekoliko godina, osobito nakon predstavljanja ChatGPT-a, alata koji je popularizirao generativnu umjetnu inteligenciju (GenAI), ona postaje sve prisutnija u svakodnevnom životu i raznim poslovnim aktivnostima. Iako umjetna inteligencija u širem smislu nije nova tema u znanstvenim istraživanjima, pojavom alata generativne umjetne inteligencije značajno se povećao interes znanstvene zajednice i šire javnosti za proučavanje prihvatanja i korištenja ovih alata (Brühl, 2023; Jovanović & Campbell, 2022). Unatoč toj popularnosti, još uvijek nedostaju teorijski utemeljeni modeli koji bi precizno objašnjavali ponašanje korisnika u ovom kontekstu, što predstavlja jasnu prazninu u znanstvenoj literaturi i potrebu za daljnjim istraživanjem.

Jedan od najčešće korištenih modela za istraživanje prihvatanja tehnologije je UTAUT2 model, predstavljen od strane Venkatesha i sur. (2012). Iako je taj model dokazao svoju primjenjivost u brojnim kontekstima (Tamilmani i sur., 2021), većina postojećih istraživanja koristi ga bez uključivanja važnih modifikatora poput dobi, spola i iskustva, ili ga primjenjuje u ograničenim kontekstima kao što je studentska populacija. To rezultira manjkom uvida u različite korisničke skupine i njihove specifične obrasce ponašanja (Tamilmani i sur., 2017; Marikyan i Papagiannidis, 2023). Postavlja se pitanje koliko je izvorni model primjenjiv na suvremene i kompleksne tehnologije i koji sve čimbenici nisu obuhvaćeni tim modelom.

Kako bi se UTAUT2 model prilagodio specifičnostima generativne umjetne inteligencije, ovaj doktorski rad predlaže proširenje modela uvođenjem dodatnih konstrukta kao što su osobna inovativnost (PI) i povjerenje (TR), koji su se u prethodnim istraživanjima pokazali kao značajni čimbenici pri prihvatanju novih tehnologija. Ujedno, rad predlaže precizno definiranje zavisne varijable „korištenje tehnologije“ (USE), čime se osigurava mogućnost njezine

primjene u širem spektru istraživačkih konteksta, jer ista ne predviđa korištenje ovakve tehnologije u izvornom radu (Venkatesh i sur., 2012). Testiranjem ovih konstrukta uz moderatore poput dobi, spola i iskustva, stvara se osnova za holističko razumijevanje ponašanja korisnika u kontekstu generativne umjetne inteligencije.

Još jedan aspekt istraživačkog problema odnosi se na geografsku dimenziju. Ovaj rad donosi usporednu analizu između Hrvatske (ekonomija u razvoju) i Ujedinjenog Kraljevstva (napredna ekonomija), oslanjajući se na IMF-ov Indeks spremnosti na umjetnu inteligenciju (IMF, 2024). Time se omogućuje analiza sličnosti i razlika u ponašanju korisnika između različitih ekonomskih i kulturnih konteksta, što doprinosi validaciji i prilagodbi proširenog UTAUT2 modela te proširuje njegovo područje primjene.

U kontekstu brzog razvoja i rastuće implementacije generativne umjetne inteligencije u poslovne, edukativne i osobne sfere, razumijevanje čimbenika koji potiču ili sputavaju njezino prihvaćanje postaje ključno ne samo za akademsku zajednicu, već i za praktičare, donositelje odluka i društvo u cjelini. Očekuje se da će generativna umjetna inteligencija transformirati tržište rada, komunikaciju i svakodnevne navike korisnika, a nepoznavanje obrazaca prihvaćanja može dovesti do digitalne isključenosti određenih skupina (Goldman Sachs, 2023; Lee, 2021). U tom smislu, jasno definiranje i adresiranje istraživačkog problema doprinosi znanosti, ali i stvara temelje za razvoj inkluzivnih, etičkih i održivih strategija implementacije ovakvih tehnologija.

1.2. Svrha i ciljevi istraživanja

Svrha doktorskog rada je na temelju relevantne literature u području prihvaćanja i korištenja tehnologije analizirati i proširiti UTAUT2 te na temelju prilagođenog modela istražiti prihvaćanje i korištenje generativne umjetne inteligencije u Republici Hrvatskoj i Ujedinjenom Kraljevstvu. Navedene zemlje odabrane su s ciljem usporedbe napredne ekonomije i ekonomije u razvoju po pitanju spremnosti na umjetnu inteligenciju, pri čemu IMF-ov Indeks spremnosti na umjetnu inteligenciju kategorizira Hrvatsku kao ekonomiju u razvoju po pitanju spremnosti na umjetnu inteligenciju, a Ujedinjeno Kraljevstvo kao naprednu ekonomiju (IMF, 2024).

Primarni cilj je na temelju teorijskih i empirijskih dokaza napraviti, testirati i proširiti konceptualni model te utvrditi utjecajne čimbenike.

Iz navedenog primarnog cilja, proizlaze i ostali specifični ciljevi koji su kategorizirani prema konceptualnom, empirijskom i aplikativnom dijelu.

Ciljevi konceptualnog dijela:

- istražiti i opisati razvoj UTAUT2 modela kroz povijest te sve njegove elemente pojedinačno
- napraviti pregled literature radova koji istražuju prihvaćanje i korištenje generativne umjetne inteligencije ili srodnih tehnologija na UTAUT2 modelu ili srodnim modelima
- prikazati dosadašnje spoznaje o prihvaćanju i korištenju generativne umjetne inteligencije
- analizirati postojeće proširene verzije UTAUT2 modela te shodno tome prilagoditi konceptualni model za istraživanje u kontekstu generativne umjetne inteligencije
- predložiti izmjene čestica zavisne varijable korištenja tehnologije koje će biti primjenjive na većem spektru istraživanja
- predložiti prilagođeni i prošireni UTAUT2 model za kontekst istraživanja generativne umjetne inteligencije.

Ciljevi empirijskog dijela:

- ispitati pouzdanost i valjanost mjernog modela te testirati utjecaje u predloženom modelu
- usporediti dobivene empirijske rezultate s rezultatima prethodnih istraživanja
- usporediti dobivene empirijske rezultate napredne ekonomije i ekonomije u razvoju po IMF-ovom Indeksu spremnosti na umjetnu inteligenciju.

Ciljevi aplikativnog dijela:

- pojasniti prediktore namjere korištenja i ponašanje pri korištenju generativne umjetne inteligencije
- predložiti odgovarajuće strategije za poticanje korištenja generativne umjetne inteligencije
- ukazati na različitosti po pitanju namjere korištenja i samog korištenja generativne umjetne inteligencije među tržišnim segmentima ovisno o njihovoj dobi, spolu ili iskustvu s navedenom tehnologijom
- predložiti načine poticanja pasivnih korisnika i nekorisnika kako bi postali aktivni korisnici alata generativne umjetne inteligencije
- predložiti preporuke za poticanje korištenja generativne umjetne inteligencije.

1.3. Hipoteze istraživanja

Hipoteze istraživanja koje slijede u ovom radu skrojene su po uzoru na one postavljene u izvornom radu Venkatesha i sur. (2012). Odnosi nezavisnih i zavisnih varijabli, kao i moderatori tih odnosa unutar hipoteza, također su preuzeti od navedenih autora i kao takvi se testiraju u ovom radu, s naglaskom na kontekst prihvatanja i korištenja generativne umjetne inteligencije.

H1: Dob i spol moderiraju utjecaj očekivanog učinka na namjeru ponašanja, tako da je utjecaj jači kod mlađih muškaraca.

H2: Dob, spol i iskustvo moderiraju utjecaj očekivanih napora na namjeru ponašanja, tako da je utjecaj jači kod mlađih žena u ranoj fazi iskustva s tehnologijom.

H3: Dob, spol i iskustvo moderiraju utjecaj društvenog utjecaja na namjeru ponašanja, tako da je utjecaj jači kod starijih žena u ranoj fazi iskustva s tehnologijom.

H4a: Dob, spol i iskustvo moderiraju utjecaj olakšavajućih uvjeta na namjeru ponašanja, tako da je utjecaj jači kod starijih žena u ranoj fazi iskustva s tehnologijom.

H4b: Dob i iskustvo moderiraju utjecaj olakšavajućih uvjeta na korištenje, tako da je utjecaj jači kod starijih, posebno onih s većom razinom iskustva s tehnologijom.

H5: Dob, spol i iskustvo moderiraju utjecaj hedonističke motivacije na namjeru ponašanja, tako da je utjecaj jači među mlađim muškarcima u ranoj fazi iskustva s tehnologijom.

H6: Dob i spol moderiraju utjecaj vrijednosti cijene na namjeru ponašanja, tako da je utjecaj jači među starijim ženama.

H7a: Dob, spol i iskustvo moderiraju utjecaj navike na namjeru ponašanja, tako da je utjecaj jači kod starijih muškaraca s većom razinom iskustva s tehnologijom.

H7b: Dob, spol i iskustvo moderiraju utjecaj navike na korištenje tehnologije, tako da je utjecaj jači kod starijih muškaraca s većom razinom iskustva s tehnologijom.

H8: Dob i spol moderiraju utjecaj osobne inovativnosti na namjeru ponašanja, tako da je utjecaj jači kod mlađih muškaraca.

H9: Dob, spol i iskustvo moderiraju utjecaj povjerenja na namjeru ponašanja, tako da je utjecaj jači među mlađim muškarcima s većom razinom iskustva s tehnologijom.

H10: Iskustvo moderira utjecaj namjere ponašanja na korištenje tehnologije, tako da je utjecaj jači kod korisnika s manje iskustva.

Istraživačko pitanje: *Je li i u kojoj mjeri moguće predvidjeti namjeru ponašanja i stvarno ponašanje korisnika generativne umjetne inteligencije na temelju predloženog konceptualnog modela koji je utemeljen na UTAUT2 modelu (očekivani učinak, očekivani naponi, društveni utjecaj, olakšavajući uvjeti, hedonistička motivacija, cijena i navika moderirani dobi, spolom i iskustvom) te proširen konstruktima osobne inovativnosti i povjerenjem?*

1.4. Metodologija istraživanja

U svrhu izrade doktorskog rada korišteni su sekundarni podatci, analizirana je relevantna znanstvena literatura koja je upotpunjena dobivenim primarnim podacima kroz empirijsko istraživanje. Prilikom prikupljanja primarnih podataka korišten je visokostrukturirani istraživački instrument napravljen prema relevantnim ljestvicama i česticama koje su prilagođene i preuzete iz dosadašnjih istraživanja (Venkatesh i sur., 2012; Venkatesh i sur., 2008; Gansser i Reich, 2021; Mohd Rahim i sur., 2008). Prikupljanje podataka provođeno je isključivo internetskim putem jer pretpostavka je kako korištenje alata generativne umjetne inteligencije podrazumijeva i pristup internetu i digitalnim kanalima komunikacije.

Empirijsko istraživanje provedeno je u dvije zemlje. Cilj primarnog istraživanja bio je istražiti prihvaćanje i korištenje generativne umjetne inteligencije u Republici Hrvatskoj, ali i po uzoru na Chopdar i sur. (2018) te Merhi i sur. (2019) napraviti komparativnu analizu s još jednom zemljom.

U kontekstu ovoga doktorskog rada to je Ujedinjeno Kraljevstvo, kako bi se izvukli određeni zaključci o sličnostima i razlikama po pitanju prediktora namjere korištenja i stvarnog korištenja generativne umjetne inteligencije kod zemalja koja su po IMF-ovom Indeksu spremnosti na umjetnu inteligenciju okarakterizirane kao ekonomija u razvoju i napredna ekonomija. Uzorak obuhvaća 400 ispitanika po zemlji, pri čemu su ispitanici u Republici Hrvatskoj pozvani na sudjelovanje tzv. metodom snježne grude (engl. *Snowball method*), koristeći komunikaciju e-poštom i društvene mreže uz zamolbu da u istraživanje pozovu članove referentne skupine iz svog okruženja. Prilikom prikupljanja podataka kontinuirano se nadzirala dobna i spolna struktura ispitanika kako bi se osigurala ravnomjernost uzorka po navedenim demografskim skupinama. Za tržište Ujedinjenog Kraljevstva koristio se Prolific, međunarodno priznata platforma za posredovanje u pronalasku ispitanika za istraživanja. Prolific nije bilo moguće koristiti za hrvatsko tržište zbog ograničenih mogućnosti selekcije ispitanika na spomenutom tržištu. Planirani uzorak ispitanika obuhvaćao je ispitanike u dobi od 18 do 65 godina različitih demografskih karakteristika, a prvo pitanje u anketi, s kojim bi

filtrirali ispitanike, bilo je o prethodnom iskustvu korištenja određenih alata generativne umjetne inteligencije. Na taj način Prolific je istraživaču omogućio da dođe samo do ispitanika koji već imaju određena iskustva s korištenjem alata generativne umjetne inteligencije, a što je iznimno važno za kontekst ovog istraživanja.

Za potrebe pisanja doktorskog rada korištene su raznovrsne metode znanstvenog istraživanja: metoda analize i sinteze, indukcije i dedukcije, apstrakcije i konkretizacije, generalizacije, deskripcije, klasifikacije, kompilacije, komparativne, povijesna, statistička metoda i druge.

Za analizu primarnih podataka korištena je univarijatna statistika (frekvencije, aritmetička sredina, standardna devijacija), multivarijatna statistika (faktorska analiza), dok je za testiranje i donošenje odluka o hipotezama korišten PLS-SEM, modeliranje strukturalnim jednadžbama kao jedna od naprednijih metoda ispitivanja utjecajnih čimbenika, ali kroz metodu parcijalnih najmanjih kvadrata (engl. *Partial Least Square*). Analiza rezultata istraživanja (univarijatne i multivarijatne statistike) izvršena je u programskom paketu *Statistical Package for the Social Science* 25 (SPSS) i JASP 0.17.2.1, a modeliranje strukturalnim jednadžbama izvršeno je metodom parcijalnih najmanjih kvadrata pri čemu korišten program SmartPLS 4.

Za potrebe ovog doktorskog rada korištena je Likertova ljestvica od 1 do 7, pri čemu su krajnje vrijednosti 1 – „u potpunosti se ne slažem“, a 7 – „u potpunosti se slažem“. Osim spomenute Likertove ljestvice i čestica konstrukta, radi usporedivosti, po uzoru na istraživanje Venkatesh i sur. (2012) preuzeta su i mjerenja za moderatore pa su tako ponuđene opcije kod spola muškarci i žene, dob je mjerena u godinama, a iskustvo u mjesecima korištenja navedene tehnologije u središtu istraživanja.

Ono što je jedna od ključnih razlika ovoga istraživanja, u metodološkom smislu, u odnosu na izvorni rad (Venkatesh i sur., 2012), jest u činjenici da je izvorni rad provodio prikupljanje podataka u dvije etape, a unutar ovoga istraživanja prikupljanje podataka bilo je u samoj jednoj etapi, ali su ispitivana dva različita tržišta.

1.5. Očekivani znanstveni doprinos

Očekivani znanstveni doprinos, u konceptualnom kontekstu, očituje se u boljem razumijevanju teorijskog okvira Ujedinjene teorije prihvatanja i korištenja tehnologije 2 (UTAUT2), analizi postojećih spoznaja te mogućih proširenja konceptualnog modela za kontekst generativne umjetne inteligencije, ali i usavršavanju zavisne varijable ponašanja korisnika / korištenja (engl. *Use behavior* / USE). Za razliku od ostalih konstrukta unutar UTAUT2 modela, zavisna varijabla stvarnog korištenja (USE) nije formulirana na način da bude univerzalno primjenjiva u različitim kontekstima istraživanja zbog čega ju je nužno prilagoditi tehnologiji u središtu istraživanja. Cilj je istraživanja napraviti konstrukt za ispitivanje stvarnog korisničkog ponašanja pri korištenju tehnologije koji će biti primjenjiv u raznim kontekstima.

U metodološkom i empirijskom kontekstu, doprinos se pak očituje u testiranju odnosa između pojedinih varijabli u proširenom konceptualnom modelu kao i u prilagođavanju i testiranju mjernih ljestvica konstrukta čime će se kreirati mjerni instrument za buduća srodna istraživanja. Primjenom PLS-SEM analize bit će ispitani izravni i neizravni utjecajni čimbenici čime će se dobiti sveobuhvatnost modela, a samim time i potencijalno veći udio objašnjenja varijance promatranog ponašanja. Nadalje, u navedenom istraživanju planiraju su koristiti sve varijable iz izvornog rada Venkatesh i sur. (2012), uključujući i moderatore koji su često neopravdano izostavljeni u istraživanjima na UTAUT2 modelu. Kako bi model bio uspješno testiran neophodno je istraživanjem obuhvatiti gotovo sve tržišne segmente u smislu dobi, spola i iskustva u korištenju generativne umjetne inteligencije, a što je također znanstveni doprinos jer su testiranja na UTAUT2 modelu prečesto usmjerena samo na tržišne skupine (najčešće studentsku populaciju). Nadalje, vidljivo je kako u metodološkom i empirijskom smislu postoji značajan istraživački prostor za istraživanja u polju generativne umjetne inteligencije jer većina radova ovu je temu tek površno istražila ispitujući samo prediktore namjere ponašanja bez da je uzela u obzir sve demografske skupine, u većini radova na ovu i srodne teme nisu korišteni moderator, a koji su jedan od najvažnijih elemenata unutar UTAUT2 modela.

Aplikativni doprinos odnosi se prvenstveno na marketinške stručnjake, stručnjake za ljudske resurse, političke i obrazovne institucije te sve poslovne subjekte čiji je cilj bolje razumijevanje potrošača u smislu korištenja generativne umjetne inteligencije. Razumijevanje čimbenika koji se planiraju istražiti unutar ovog doktorskog rada pomoći će marketinškim stručnjacima da usmjere svoje napore na elemente koji snažno utječu na namjeru korištenja i stvarno korištenje alata generativne umjetne inteligencije i srodnih tehnologija. Proširivanje znanstveno i stručno

prepoznatog UTAUT2 modela pruža opsežniji i precizniji okvir za razumijevanje čimbenika usvajanja spomenutih tehnologija među različitim tržišnim skupinama. Ovakve spoznaje mogu biti značajne za tvrtke u raznim sektorima, ali posebno za one koje u uporabi generativne umjetne inteligencije vide priliku i interes za vlastitim napretkom. Očekuje se kako će rezultati ovog istraživanja doprinijeti boljem razumijevanju usvajanja i korištenja generativne umjetne inteligencije među članovima svih demografskih skupina u Republici Hrvatskoj, a čiji će se rezultati usporediti s korisnicima generativne umjetne inteligencije iz Ujedinjenog Kraljevstva. Usmjeravanjem marketinških strategija na identificirane prediktore u namjeri korištenja i korištenju generativne umjetne inteligencije, marketinški stručnjaci mogu optimizirati integraciju navedene tehnologije, poboljšati korisničko iskustvo i steći konkurentsku prednost u iznimno dinamičnom okruženju.

1.6. Struktura rada

Struktura doktorskog rada osmišljena je na način da prati logički slijed znanstvenog istraživanja, od teorijske osnove i analize prethodnih istraživanja, preko razvoja modela pa sve do provedbe i analize empirijskog istraživanja. Rad je podijeljen na devet glavnih poglavlja, a započinje s uvodom u kojemu se definira kontekst istraživanja, objašnjava problematika i navode ciljevi rada. Unutar ovog dijela obrađene su i hipoteze te istraživačko pitanje, metodologija te očekivani znanstveni i aplikativni doprinosi doktorskog rada. Uvodom se čitatelju pruža jasan okvir za razumijevanje sadržaja doktorskog rada.

U nastavku, drugo poglavlje daje teorijsko uporište i definira umjetnu inteligenciju. Objašnjavaju se pojmovi umjetne inteligencije s naglaskom na generativnu umjetnu inteligenciju. Detaljno su predstavljene evolucijske faze umjetne inteligencije, definicije te etički aspekti upotrebe ovako napredne tehnologije.

Treće poglavlje doktorskog rada naglasak stavlja na modele prihvaćanja tehnologije. Naglasak je stavljen na razvoj modela UTAUT i UTAUT2, njihovu evoluciju, teorijska polazišta i ključne konstrukte i odnose unutar modela koji su srž empirijskog istraživanja ovog doktorskog rada.

U četvrtom poglavlju prikazan je pregled literature. Analizirane su dosadašnje studije koje se bave prihvaćanjem i korištenjem generativne umjetne inteligencije i srodne tehnologije mjerene UTAUT2 modelom ili srodnim modelima. Pritom su identificirani konstrukti koji će biti dodatno uključeni u prošireni model ovog rada – osobna inovativnost i povjerenje. Osim toga, četvrto poglavlje posvećeno je i analizi postojećih znanstvenih spoznaja iz ovog istraživačkog

područja. Uključena je i analiza preliminarnog istraživanja koje je pomoglo u razvoju konačnog mjernog instrumenta.

U petom poglavlju obrađena je usporedba između Ujedinjenog Kraljevstva i Republike Hrvatske u kontekstu njihove spremnosti na umjetnu inteligenciju. Navedena usporedba daje kontekst boljeg razumijevanja rezultata empirijskog istraživanja.

Šesto poglavlje opisuje konceptualni dizajn istraživanja, uključujući i razvoj istraživačkog instrumenta, proširenja i modifikacije UTAUT2 modela te definiranje hipoteza i istraživačkog pitanja.

U sedmom poglavlju prikazani su rezultati empirijskog istraživanja te je detaljno objašnjen proces prikupljanja podataka, opis uzorka, analiza rezultata i testiranje hipoteza istraživanja. Hipoteze su testirane na cjelovitom uzorku, a potom su rezultati analizirani zasebno za tržište Ujedinjenog Kraljevstva i za tržište Republike Hrvatske.

U osmom poglavlju koje se bavi raspravom, interpretiraju se rezultati istraživanja te se dobiveni rezultati uspoređuju s postojećim znanstvenim spoznajama iz ovog istraživačkog područja. Nadalje, raspravljaju se implikacije istraživanja, ograničenja rada te su napravljene preporuke za buduća istraživanja.

U konačnici, deveto poglavlje rada donosi zaključke doktorskog rada u kojem se sažimaju ključni nalazi i doprinosi ovog doktorskog rada te se nude završne misli autora. Nakon zaključka, dolaze popisi literature, tablica, slika, grafikona, kratica i druga popratna dokumentacija.

Ovakvom strukturom rada omogućeno je sustavno praćenje cijelog znanstveno-istraživačkog procesa, od definiranja istraživačkog problema do krajnjih zaključaka čime je osigurana znanstvena relevantnost i metodološka utemeljenost doktorskog rada.

2. UMJETNA INTELIGENCIJA

Umjetna inteligencija najfascinantnija je tehnologija današnjice te vrlo popularna tema u mnogim znanstvenim i stručnim raspravama, a razlog tomu leži u činjenici što ima sposobnost imitirati ljudsku inteligenciju (Mondal, 2020). Kako bi se precizno definirao pojam umjetne inteligencije, potrebno je detaljno razmotriti što to uopće predstavlja pojam „inteligencija“. Cambridgeov rječnik (*Cambridge Dictionary*, n.d.) inteligenciju opisuje kao sposobnost učenja, razumijevanja, prosuđivanja i stvaranja mišljenja temeljenog na razumu. Inteligencija, koja dolazi od latinske riječi *intelligentia*, označava razum, vještinu i razboritost, sposobnost mišljenja koja omogućava snalaženje u novim i nepoznatim situacijama (Enciklopedija, n.d.).

Flasinski (2016) navodi kako su veliki umovi poput Aristotela, Tome Akvinskog, Williama Occama, Renea Descartesa, Thomasa Hobbesa i Leibniza propitivali o mogućnostima automatizacije rasuđivanja, o potpunom razumijevanju kognitivnih operacija. Preteča je to umjetne inteligencije koja nije bila moguća sve do sredine 20. stoljeća i pojave prvih računala, kada se doista kreće realno promišljati o strojevima koji izvršavaju određene operacije imitirajući rad ljudskoga uma.

Indeks umjetne inteligencije, koji Sveučilište Stanford svake godine objavljuje u formi izvješća trenutnog stanja, trendova i percepcije u vezi umjetne inteligencije te njenog utjecaja na društvo, jedan je od najrelevantnijih javno dostupnih dokumenata ovoga tipa. Unutar navedenog izvješća navodi se deset ključnih stavki koje definiraju umjetnu inteligenciju i njezine mogućnosti do 2024. godine (Maslej i sur., 2024):

1. *Umjetna inteligencija nadmašuje ljude u nekim zadacima, ali ne svim zadacima.*
Umjetna inteligencija nadmašila je izvedbu čovjeka u nekoliko zadataka kao što su prepoznavanje slika, vizualno zaključivanje ili razumijevanje engleskog jezika. Ipak, umjetna inteligencija zaostaje u složenijim zadacima poput matematičkih problema na natjecateljskoj razini, zdravorazumskog razumijevanja i planiranja.
2. *Industrija i dalje dominira istraživanjem na području napredne umjetne inteligencije.*
U 2023. godini, industrija je proizvela 51 napredni model strojnog učenja dok je akademska zajednica doprinijela sa samo 15 modela. Ipak, 21 model nastao je kroz suradnju industrije i akademije, što je dosadašnji rekord.
3. *Napredni modeli postaju znatno skuplji.* Troškovi treniranja najmodernijih AI modela dosegli su neviđene razine. Primjerice, treniranje naprednog modela kao što je GPT-

4 kojeg koristi OpenAI procjenjuje se na 78 milijuna dolara računalnih resursa, dok je Googleov Gemini Ultra koštao 191 milijun dolara.

4. *Sjedinjene Američke Države vode u utrci protiv Kine, Europske Unije i Ujedinjenog Kraljevstva kao vodeći izvor najboljih modela umjetne inteligencije.* U 2023. godini, 61 značajan model umjetne inteligencije potekao je iz američke institucija, nadmašivši Europsku uniju koja je proizvela 21 i Kinu koja je proizvela 15 modela.
5. *Nedostatak robusnih i standardiziranih procjena za odgovornost velikih jezičnih modela ozbiljan su problem.* Indeks umjetne inteligencije otkriva značajan nedostatak standardizacije u odgovornom izvještavanju o umjetnoj inteligenciji. Vodeći kreatori velikih jezičnih modela, uključujući OpenAI, Google i Anthropic, uglavnom testiraju svoje modele na različitim odgovornim mjerilima za umjetnu inteligenciju što značajno otežava sustavno uspoređivanje rizika i ograničenja vrhunskih modela umjetne inteligencije.
6. *Investicije u generativnu umjetnu inteligenciju naglo rastu.* Unatoč opadanju ukupnih privatnih ulaganja u umjetnu inteligenciju, financiranje generativne umjetne inteligencije gotovo se udvostručilo u odnosu na 2022. godinu, dosegnuvši 25,2 milijarde američkih dolara.
7. *Podatci su jasni – umjetna inteligencija radnike čini produktivnijima i vodi ka kvalitetnijem radu.* U 2023. godini, nekoliko studija procijenilo je utjecaj umjetne inteligencije na radnu snagu, sugerirajući da umjetna inteligencija omogućuje radnicima da brže dovrše zadatke i poboljšaju kvalitetu svoga rada. Nadalje, studije pokazuju kako umjetna inteligencija utječe na smanjenje jaza u vještinama radnika s nižim i višim kvalifikacijama. Ipak, određene studije ukazuju i na to kako korištenje umjetne inteligencije bez odgovarajućeg nadzora vodi ka smanjenju učinkovitosti.
8. *Zahvaljujući umjetnoj inteligenciji, znanstveni napredak se značajno ubrzava.* U 2022. godini umjetna inteligencija značajno je počela unapređivati znanstvena otkrića, dok je 2023. godine donijela lansiranje još mnogo značajnih aplikacija povezanih sa znanošću.
9. *Broj propisa vezanih za umjetnu inteligenciju u SAD-u naglo raste.* Broj propisa vezanih za umjetnu inteligenciju tijekom 2023. godine u Sjedinjenim Američkim Državama značajno je porastao. U 2023. godini bilo je 25 AI-povezanih propisa, a 2016. godine bilo je samo jedan. Samo tijekom 2023. godine broj AI-povezanih propisa narastao je za 56,3 %.
10. *Ljudi diljem svijeta sve su svjesniji potencijalnog utjecaja umjetne inteligencije, ali i sve zabrinutiji.* Ankete Ipsosa i Pew Researcha ukazuju na činjenicu da se tijekom 2023.

godine povećao postotak onih koji misle da će AI dramatično utjecati na njihove živote u narednih tri do pet godina. Ipak, 52 % ispitanika izrazilo je zabrinutost zbog AI proizvoda i usluga, što je veliki rast u odnosu na istraživanje provedeno 2022. godine.

Kai-Fu Lee (2021), jedan od vodećih svjetskih stručnjaka za umjetnu inteligenciju, za magazin Time napisao je kako bi umjetna inteligencija mogla biti tehnologija koja je napravila najveću promjenu u povijesti čovječanstva. Nadalje, Kai-Fu Lee (2021) dodaje kako umjetna inteligencija, kao i većina drugih tehnologija, nije ni dobra niti zla, ali da će u konačnici čovječanstvu donijeti više pozitivnih nego negativnih učinaka. U knjizi „AI 2041“, Kai-Fu Lee i Chen Qiufan (2021) predviđaju budućnost u suživotu s umjetnom inteligenciju te najavljuju kako će navedena tehnologija u naše privatne i poslovne živote ulaziti kao asistent, podržavajući alat, ali da će određene rutinske poslove s vremenom u potpunosti zamijeniti. Nadalje, najavljuju revoluciju u zdravstvu, prometu, obrazovanju te raznim drugim sferama privatnog i poslovnog života.

2.1. Razvoj umjetne inteligencije

Umjetna (od čovjeka stvorena) inteligencija (sposobnost razmišljanja) odnosi se na strojeve, uređaje, računala koji imaju funkciju da donose odluke i djeluju simulirajući ljudsku inteligenciju zbog čega ih se zove inteligentnim uređajima. Ipak, da bismo mogli zaključiti kako je uređaj inteligentan on mora moći učiti, zaključivati i na temelju naučenog donositi odluke. Drugim riječima, uređaj ili stroj možemo nazvati inteligentnim ukoliko je u stanju proći Turingov test (Mondal, 2020).

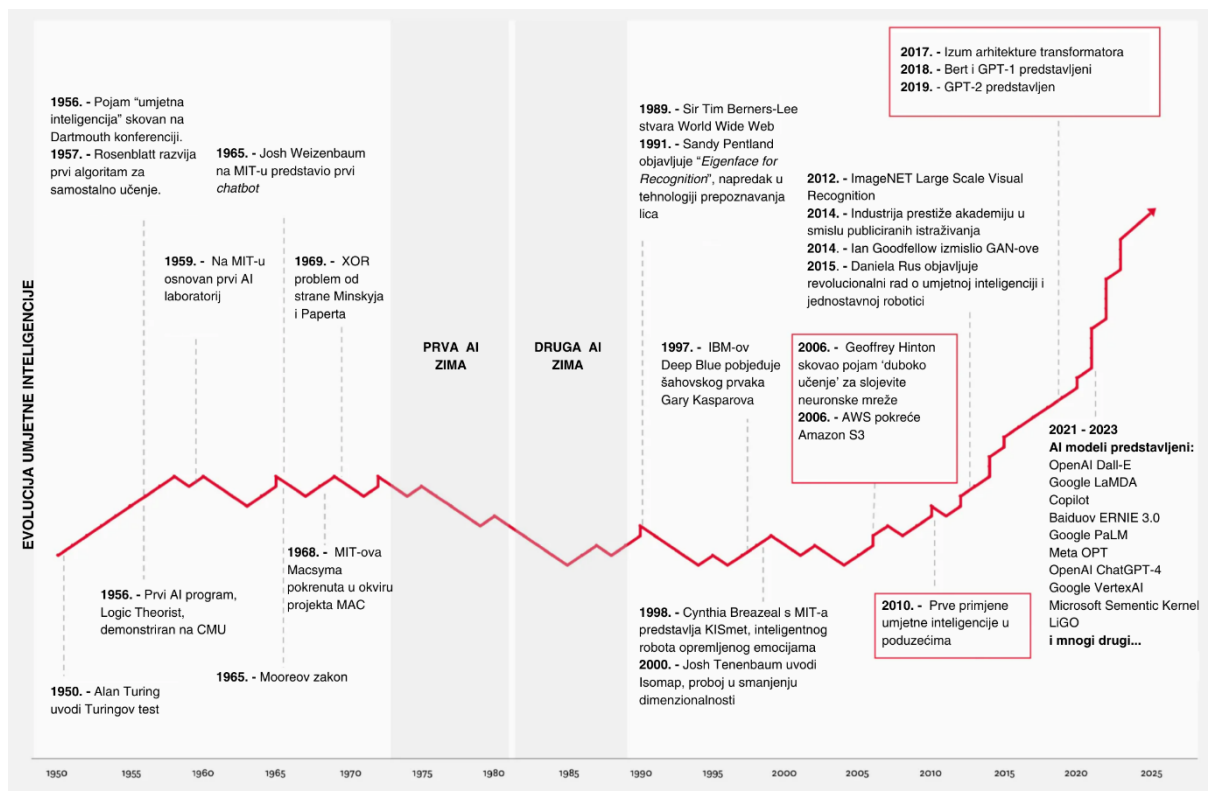
Turingov test predstavio je Alan Turing, često oslovljavan kao otac računarstva i umjetne inteligencije (Sharkley, 2012), a navedenim je testom, prije više od 70 godina, postavio temelje za razvoj umjetne inteligencije. Turing (1950) je rad pod nazivom „Računalni strojevi i inteligencija“ (engl. *Computing machinery and intelligence*) započeo pitanjem „Mogu li strojevi razmišljati / misliti?“ (engl. *Can machines think?*). Spomenuti test izvornog naziva „Igra oponašanja“ (engl. *Imitation game*) danas je poznat kao Turingov test i osnova je gotovo svake edukacije o umjetnoj inteligenciji. U načelu, Turingov test mjeri sposobnost stroja da pokaže inteligentno ponašanje i ukoliko evaluator ne može procijeniti komunicira li sa strojem ili s čovjekom, smatra se da je stroj prošao Turingov test (Seseri i sur., 2024).

Mondal (2020) predstavlja vremensku liniju važnih događaja za razvoj umjetne inteligencije:

- 1837. – prvi model programabilnog računala; Charles Babbage, Ada Lovelace

- 1943. – predstavljena neuronska mreža; McCulloch i Walter Pitts
- 1950. – Turingov test: kako testirati inteligenciju stroja; Alan Turing
- 1955. – skrojen pojam „Umjetna inteligencija (AI)“; John McCarthy (Dartmouth)
- 1965. – prvi chatbot ELIZA; MIT – laboratorij za razvoj umjetne inteligencije
- 1980. – ekspertni sustav; Edward Feigenbaum
- 1980. – neuronska mreža korištena za autonomna vozila; C. M. U.
- 1997. – Deep Blue pobjeđuje svjetskog prvaka u šahu Garryja Kasparova
- 2002. – Roomba, autonomni robot za čišćenje prostora; iRobot
- 2009. – prvi samovozeći automobil; Google
- 2011. – Watson pobjeđuje prvake u američkom kvizu Jeopardy; IBM
- 2011. - 2014. – Siri (Apple), Google Now, Cortana (Microsoft)
- 2014. – generativna protivnička mreža (GAN); Ian Goodfellow
- 2017. – Alpha Go pobjeđuje svjetskog prvaka u igri Go, Ke Jie.

S obzirom na to da je Mondalova knjiga objavljena 2020. godine, u njoj nedostaje vrlo bitan period važan za razvoj umjetne inteligencije koji se dogodio nakon toga. Kralj i sur. (2024) u svojoj verziji navode kako je od 2020. godine do 2024. godine predstavljen DALL-E (alat generativne umjetne inteligencije za generiranje vizualnih sadržaja), napisan je i zakon o umjetnoj inteligenciji koji je 2021. godine predstavila Europska unija. Samo godinu kasnije, u studenom 2022. godine, javnosti je predstavljen ChatGPT, alat generativne umjetne inteligencije, odnosno veliki jezični model koji je u manje od dva mjeseca skupio 100 milijuna korisnika i masovno popularizirao upotrebu generativne umjetne inteligencije. Dakako, u knjizi nije prisutan ni DeepSeek, kineski AI model koji je uzdrmao tržište i prikupio više od 33 milijuna aktivnih korisnika u prvih 30 dana od lansiranja (Jiang, 2025).



Slika 1. Prikaz bitnih događaja u razvoju umjetne inteligencije

Izvor: prilagođeno prema: Seseri i sur. (2024)

Kako je vidljivo na slici 1., jedan od najvažnijih događaja u ranim fazama razvoja umjetne inteligencije je napredak u razvoju mikročipova. Mooreov zakon predstavljen 1965., nazvan po njegovom začetniku Gordonu Mooreu, predviđa kako se broj tranzistora na mikročipu udvostručuje otprilike svake dvije godine što rezultira eksponencijalnim rastom računalne snage i učinkovitosti uz smanjenje troškova (Seseri i sur., 2024; Moore, 1965). Umjetna inteligencija doživjela je nekoliko hladnih razdoblja u kojima se nije događalo puno po pitanju razvoja umjetne inteligencije te su navedena doba u literaturi poznata kao „AI zime“ (engl. *AI winters*). Umjetna inteligencija podložna je takvim ciklusima upravo zbog prevelikih očekivanja od takve tehnologije (Floridi, 2020). Zimska razdoblja umjetne inteligencije tijekom 70-ih i 80-ih godina prošlog stoljeća potaknuli su prijelaz s pretjerano optimističnih predviđanja na pragmatičniji pristup razvoju ove tehnologije. U vremenima smanjenog interesa i financiranja potiče se razvoj učinkovitijih i isplativijih rješenja, a takav je slučaj bio i s umjetnom inteligencijom te je tijekom 1990-ih godina ponovno doživjela snažan rast (Seseri i sur., 2024).

Sam naziv umjetna inteligencija postao je breme za istraživače i tvrtke koje su se bavile razvijanjem takve tehnologije pa su se tijekom 90-ih godina prošlog stoljeća napredna rješenja

stavljala pod nazivnik napredne informatike i računalne inteligencije. Tijekom 1990-ih umjetna inteligencija nije dobila puno zasluga premda je bila integrirana u brojna softverska rješenja, aplikacije i algoritme (Markoff, 2005; NRC, 1999; Manyinka, 2022; McCorduck & Cfe, 2004).

Sredinom 1990-ih, točnije 1996. godine, dogodio se jedan od najbitnijih događaja u popularizaciji umjetne inteligenciji i njenih mogućnosti – IBM-ov Deep Blue porazio je Garyja Kasparova, šahovskog velemajstora i svjetskog prvaka. Premda je Deep Blue u prvom pokušaju Kasparova porazio u samo dvije od ukupno šest rundi, bilo je to dovoljno da umjetna inteligencija fascinira cijeli svijet jer je u šahu, igri koja predstavlja intelektualno nadmetanje, barem djelomično porazila najboljeg čovjeka na svijetu. Već iduće godine, 1997. godine, Deep Blue potpuno je porazio Kasparova s 3,5 – 2,5 (Krajcar, n.d.; IBM, 2020).

Najveći napredak u razvoju umjetne inteligencije pri kraju 20. stoljeća bio je okretanje ka sofisticiranim matematičkim pristupima kao što su neuronske mreže, probabilističko zaključivanje, statistika i teorija odlučivanja. Oslanjajući se na matematičke modele, umjetna inteligencija pronasla je put ka integraciji u raznim sustavima kao što su videoigre, sustavi za rezervaciju, automatizirani prevoditelji, algoritmi za tražilice i mnogi drugi (McCorduck i Cfe, 2004).

Početkom 21. stoljeća započeo je i novi evolucijski val obilježen eksponencijalnim rastom podataka generiranih digitalnim tehnologijama, ali i napretkom računalne snage što je otvorilo ogromne mogućnosti napretku umjetne inteligencije. Integracija velikih skupova podataka u algoritme strojnog učenja omogućila je razvoj sofisticiranijih i preciznijih prediktivnih modela (Russell i Norvig, 2021). Nakon strojnog učenja došlo je i duboko učenje koje koristi više slojeva jednostavnih i prilagodljivih računalnih elemenata. Eksperimenti s dubokim učenjem počeli su još 1970-ih godina, no tek 2011. godine duboko učenje pokazuje značajnije rezultate, prvenstveno u području prepoznavanja govora i prepoznavanja vizualnih predmeta (Russell i Norvig, 2021). Drugo desetljeće 21. stoljeća obilježeno je tehnološkim inovacijama poput pojave virtualnih asistenata kao što su Siri i Alexa te drugih srodnih *smarthome*¹ tehnologija, pojave Teslinog autopilota, ali i prvih sofisticiranih *deepfake*² uradaka napravljenih uz pomoć umjetne inteligencije (Tyson, 2022). Ipak, velike stvari tiho su se razvijale u pozadini, referirajući se na razvoj velikih jezičnih modela i srodnih alata generativne umjetne

¹ *Smarthome* tehnologija podrazumijeva sustav povezanih uređaja i sustava unutar doma koji omogućuju automatizirano i daljinsko upravljanje funkcijama poput rasvjete, grijanja, hlađenja, sigurnosti ili zabave.

² *Deepfake* je tehnologija koja koristi umjetnu inteligenciju za stvaranje realističnih, ali lažnih videozapisa, audiozapisa ili slika, najčešće tako što se lice ili glas jedne osobe zamijeni licem ili glasom druge osobe.

inteligencije koji će svjetlo dana ugledati nakon gotovo desetljeća razvoja (Wang, 2024; Mueller, 2024).

U znanstvenom radu pod naslovom „*Attention is All You Need*“ (Vaswani i sur., 2017), osam Googleovih istraživača predstavilo je transformerski model, zahvaljujući kojem ih danas nazivaju tvorcima moderne umjetne inteligencije (Levy, 2024). Smatra se to jednim od najvećih znanstvenih okrića u području razvoja umjetne inteligencije jer je značajno unaprijedilo rukovanje sekvencijalnim podacima u umjetnoj inteligenciji, zamijenivši dotadašnje dominantne rekurentne neuronske mreže (RNN). Uvođenjem mehanizma pažnje, transformerski model učinkovito je riješio ograničenja RNN-ova, poput poteškoća s učenjem dugoročnih ovisnosti i računalne neefikasnosti (Seseri i sur., 2024).

Od 2020. godine događa se veliki „*AI boom*“ potaknut s nekoliko ključnih otkrića u dizajniranju neuronskih mreža, sve većom dostupnošću podataka, hardverskim napretkom, ali prvenstveno zbog spremnosti velikih tehnoloških tvrtki da plate ogromne razine računalne snage (Chow i Perrigo, 2023). Stanford ovaj rapidni razvoj događaja u svijetu umjetne inteligencije naziva „AI proljeće“ (Bommasani, 2023). Statista (2024) navodi kako je procjena veličine svjetskog tržišta umjetne inteligencije danas 244 milijarde američkih dolara, a procjena je kako će veličina tog tržišta u 2030. godini biti otprilike 826 milijardi američkih dolara. S druge strane, procijenjeni rast tržišta samo generativne umjetne inteligencije penje se s 5,51 milijardi američkih dolara iz 2020. godini na 356 milijardi američkih dolara u 2030. godini.

U veljači 2020. godine Microsoft je predstavio projekt Turing – NLG, jezični model od 17 milijardi parametara (Sterling, 2020). Ipak, to nije privuklo veliki interes javnosti kao kada je OpenAI lansirao ChatGPT u studenom 2022. godine. Navedeni alat generativne umjetne inteligencije privukao je toliki interes javnosti da je već u siječnju 2023. godine imao više od 100 milijuna aktivnih korisnika. Bio je to najbrži rast neke tehnologije ili aplikacije ikada. Usporedbe radi, TikToku je trebalo devet mjeseci, a Instagramu dvije i pol godine dok su dosegli brojku od 100 milijuna korisnika (Hu, 2023). Naravno, ChatGPT-3 nije bio savršen te je usprkos općom fasciniranošću šire javnosti imao i dosta nedostataka zbog kojih je bio na udaru kritika, u prvom redu zbog haluciniranja, ali i prigovora oko krađe autorskih prava (Reed, 2024; Rachini, 2022).

Već u ožujku 2023. godine OpenAI lansirao je i ChatGPT-4, alat koji je postao više od velikog jezičnog modela, postao je multimodalan. ChatGPT-4 došao je kao unaprijeđena verzija s većim mogućnostima nego prethodnik, ali je došao i s pretplatom od 20 američkih dolara (Weitzman,

2023). Shvativši da je OpenAI već pokupio slavu s izbacivanjem prvog javno dostupnog velikog jezičnog modela, Google je požurio s lansiranjem Barda, rivala ChatGPT-u, koji nije previše poživio prije nego li je Google ugasio projekt Bard, pretvorivši ga u projekt Gemini već u prosince iste godine (Pichai, 2023; Pierce, 2023).

U prvoj polovici 2023. godine događaju se velike promjene u svijetu umjetne inteligencije, velike tehnološke kompanije konstantno su najavljivale još bolje i veće modele što je rezultiralo otvorenim pismom pod naslovom „Zaustavite divovske AI eksperimente“ (engl. *Pause Giant AI Experiments: Open Letter*), koji je potpisalo više od 30 tisuća ljudi, a među njima i Yoshua Bengio, Stuart Russell, Elon Musk, Steve Wozniak, Emad Mosaque, Andrew Yang, John Hopfield i još stotine drugih stručnjaka za umjetnu inteligenciju. U navedenom pismu stoji: „Trebamo li dopustiti strojevima da preplave naše informacijske kanale propagandom i neistinom? Trebamo li automatizirati sva radna mjesta, uključujući i ona ispunjavajuća? Trebamo li razviti nečovječne umove koji bi nas mogli brojčano nadmašiti, nadmudriti i zamijeniti? Trebamo li riskirati gubitak kontrole nad našom civilizacijom? Takve odluke ne smiju biti prepuštene neizabranim tehnološkim liderima. Moćni AI sustavi trebali bi se razvijati tek kada budemo sigurni da će njihovi učinci biti pozitivni, a rizici upravljivi. (...) Čovječanstvo može uživati u naprednoj budućnosti s umjetnom inteligencijom. Nakon uspješnog stvaranja moćnih AI sustava, sada možemo uživati u „ljetu AI-a“ u kojem ćemo ubirati plodove, projektirati te sustave za jasnu dobrobit svih i dati društvu priliku da se prilagodi. Društvo je već zaustavilo druge tehnologije s potencijalno katastrofalnim učincima na društvo. To možemo učiniti i ovdje. Uživajmo u dugom AI ljetu, a ne srljajmo nepripremljeni u jesen“ (Future of Life, 2023).

Kao odgovor na istraženu zabrinutost istaknutih stručnjaka, a shodno tome i šire javnosti, u rujnu 2023. godine Senat Sjedinjenih Američkih Država saziva Forum o uvidu u umjetnu inteligenciju gdje okupljaju vodeće tehnološke lidere SAD-a s ciljem dijaloga o benefitima i opasnostima koje umjetna inteligencija donosi društvu (Johnson, 2023). Samo mjesec dana kasnije, predsjednik Sjedinjenih Američkih Država potpisao je Izvršnu naredbu o sigurnom, zaštićenom i pouzdanom razvoju i korištenju umjetne inteligencije (The White House, 2023). Riječ je o iznimno važnom dokumentu, uz Akt o umjetnoj inteligenciji Europske unije (Europski parlament, 2024a), jednom od najbitnijih pravnih okvira za razvoj i korištenje umjetne inteligencije.

Kako je ranije spomenuto, umjetna inteligencija nije niti zla niti dobra, no sa sobom nosi određene benefite, ali i određene rizike, kao i većina drugih tehnologija (Lee, 2021). Ipak, riječ

je o vrlo softisciranoj tehnologiji koja imitira ljudsku inteligenciju i posjeduje ogromnu bazu znanja zbog čega će pomoći ljudima u još većim ostvarenjima. Gotovo sve što je civilizacija do sada stvorila plod je ljudske inteligencije, no u ovom trenutku civilizacija je na prekretnici, od ovoga trenutka strojna će inteligencija proizvoditi, osmišljavati, ubrzavati i otkrivati. Kako je sugerirao Demis Hassabis, nobelovac i direktor Googleovog DeepMind: „Prvo riješite umjetnu inteligenciju, a potom koristite umjetnu inteligenciju za rješavanje svega ostalog“. Ipak, prije nego li se „riješi“ umjetna inteligencija, potrebno je temeljito razmotriti rizike i prijetnje upotrebe i zloupotrebe takve moćne tehnologije (Russell i Norvig, 2021). Umjetna inteligencija u ovom se trenutku razvija iznimno brzo. Kako je ranije spomenuto, Gordon Moore, suosnivač Intela, 1965. godine iznio je opažanje koje kaže kako se broj tranzistora u mikročipovima udvostručuje svake dvije godine. Kroz desetljeća nakon toga taj se eksponencijalni rast računalne snage i smanjenja relativnih troškova nastavio te je navedeno opažanje postalo poznato kao Mooreov zakon (Moore, 1965). Nasuprot tomu, alati generativne umjetne inteligencije doživljaju još brži rast, istraživanje njihovih kapaciteta pokazuje kako se mogućnosti alata generativne umjetne inteligencije da rješavaju kompleksne zadatke udvostručuju svakih sedam mjeseci (Kwa i sur., 2025). Upravo je to pokazatelj kako je ovo vrlo bitan trenutak za postavljanje regulatornih, zakonskih i moralnih pravila za upravljanje ovako moćnom tehnologijom, prije nego li postane prekasno.

Ujedinjeni narodi kao jednu od glavnih prijetnji po pitanju upotrebe umjetne inteligencije vide kod korištenja smrtonosnog autonomnog oružja, a čija definicija podrazumijeva oružje koje može locirati, odabrati i eliminirati ljudske ciljeve bez ljudske intervencije. Tehnologije potrebne za takvo oružje slične su onima za razvoj autonomnih vozila i zbog toga se UN ovim pitanjem bavi već duže od desetljeća. Umjetna inteligencija u tom smislu može imati i pozitivnu i negativno ulogu, ali zbog straha od zloupotrebe takvih sustava UN se bori za izglasavanje zakona koji bi to zabranio. Glavni tajnik UN-a ponovio je poziv međunarodnoj zajednici s ciljem da se do 2026. godine nađe konsenzus za izglasavanje ovakvog zakona. (UN, 2024; Russell & Norvig, 2021).

Korištenje umjetne inteligencije za nadzor i uvjeravanje također izaziva veliku zabrinutost kod svih koji se bave etičkim pitanjem upotrebe ovakvih naprednih tehnologija. Nadziranje telefonskih linija, e-pošte, video snimki i drugih komunikacijskih kanala skupa je, zamorna i nerijetko pravno upitna aktivnost, no uz umjetnu inteligenciju sve to postaje puno brže, jeftinije, jednostavnije i masovnije. Jednako tako, upotreba umjetne inteligencije u svrhu manipuliranja i kontroliranja informacija i dezinformacija na društvenim mrežama i srodnim platformama

velika je briga za cjelokupno društvo, a zabrinutost je postala vrlo očita još 2016. godine na izborima u SAD-u (Russell & Norvig, 2016).

Pristranost umjetne inteligencije jedna je od temeljnih kritika sustavima pogonjenima umjetnom inteligencijom. Sustavi umjetne inteligencije nerijetko daju pristrane rezultate jer su, naučeni na ogromnim bazama podataka, rezultate prilagodili analitičkim podacima i zbog toga mogu diskriminirati razne društvene skupine na temelju njihova spola, rase ili po nekim drugim kriterijima. Primjerice, banka Goldman Sachs dospjela je u centar medijske pozornosti jer je utvrđeno kako je njihov algoritam navodno diskriminirao žene odobravajući im manje kreditne limite nego li muškarcima (Blackman, 2020).

Od svih strahova povezanih s umjetnom inteligencijom, onaj da će ljudi dobivati otkaze i biti zamijenjeni umjetnom inteligencijom vjerojatno je najveći. Taj strah nije neutemeljen jer podaci znanstvenih istraživanja potvrđuju činjenicu da je umjetna inteligencija sve sposobnija u obavljanju zadataka koje su do sada radili isključivo ljudi i neminovno je kako će rezultat toga biti promjene na tržištu rada (Huang & Rust, 2018). Goldman Sachs (2023) procjenjuje kako je otprilike dvije trećine zanimanja u SAD-u izloženo određenom stupnju automatizacije, ali i dodaju kako to ne znači da su sva ta zanimanja u riziku od otkaza već da će im umjetna inteligencija pomoći da budu produktivniji u svom radu. Drugim riječima, umjetna inteligencija možda neće tako brzo i efikasno zamijeniti ljude, ali ljudi koji koriste navedenu tehnologiju vjerojatno hoće one koji ju ne koriste (Lakhani, 2023). Nadalje, Goldman Sachs (2023) navodi kako 60 % radnika danas radi poslove koji nisu postojali prije 1940. godine, što implicira da se više od 85 % rasta zaposlenosti u zadnjih 80-ak godina može objasniti tehnološkim razvojem te da će shodno tome i razvoj umjetne inteligencije stvoriti prostor za neka nova radna mjesta.

Još jedan od opravdanih strahova povezanih s umjetnom inteligencije je onaj o sigurnosnim prijetnjama. Akteri koji zloupotrebljuju umjetnu inteligenciju koriste ju za pisanje malicioznih softvera, skripti za hakiranje, *phishing*, *deepfake* i slične aktivnosti. S druge strane, umjetna inteligencija postaje i važan saveznik stručnjacima za kibernetičku sigurnost, pogotovo u domeni prikupljanja podataka i dokaza, analizi istih, procjeni rizika, reviziji i brojnim drugim aktivnostima koje će transformirati profesiju kibernetičke sigurnosti u narednim godinama (Dimitriadis, 2024).

Još jedan od problema koji se često pojavljuje kod uporabe umjetne inteligencije i izaziva velike probleme je problem haluciniranja modela, tzv. halucinacije umjetne inteligencije, također poznate i kao konfabulacije umjetne inteligencije (Jamaluddin i sur., 2023). Riječ je o

rezultatima generativne umjetne inteligencije koji nisu činjenično točni, ne odgovaraju podacima iz stvarnog svijeta i često su potpuno izmišljeni. Problem halucinacija jedan je od najznačajnijih izazova jer dovodi u pitanje vjerodostojnost svega što generiraju alati generativne umjetne inteligencije (Jamaluddin i sur., 2023). Posljedice halucinacija mogu biti vrlo opasne, primjerice - medicinski sustav umjetne inteligencije zbog svog bi haluciniranja mogao netočno identificirati benignu leziju kože kao zloćudnu, što bi rezultiralo nepotrebnim medicinskim intervencijama na koži (IBM, 2024).

2.2. Definiranje umjetne inteligencije

U ljeto 1965. godine, na poziv John McCarthyja, okupilo se desetak istaknutih znanstvenika koji su imali interes za neuronske mreže, teoriju automata i proučavanje inteligencije u sklopu šestotjednog studija pod nazivom Dartmouth College Project. Navedeni događaj često se smatra začetkom umjetne inteligencije jer je upravo tada i prvi put upotrijebljen naziv „Umjetna inteligencija“ (engl. *Artificial intelligence*). Želja da se krene dublje istraživati ideja o umjetnoj inteligenciji prenesena je i na papir kada su okupljeni znanstvenici napisali prijedlog podnesen zakladi Rockefeller koja je osiguravala financijska sredstva za događaj (Bostrom, 2014). U tekstu je stajalo: „Predlažemo da se provede dvomjesečna studija umjetne inteligencije s deset znanstvenika. (...) Studija će se temeljiti na pretpostavci da se svaki aspekt učenja ili bilo koja druga značajka inteligencije može u načelu tako precizno opisati da se može simulirati pomoću stroja. Napraviti će se pokušaj otkrivanja načina na koji strojevi mogu koristiti jezik, oblikovati apstrakcije i koncepte, rješavati vrste problema koje su do sada bile rezervirane za ljude te poboljšati sami sebe. Vjerujemo da se može postići značajan napredak na jednom ili više ovih problema ako pažljivo odabrana skupina znanstvenika radi na tome tijekom ljeta“ (McCarthy i sur., 1955). Od tog velikog događaja za ovu znanstvenu disciplinu pa sve do danas, područje umjetne inteligencije prolazilo je različita razdoblja entuzijazma i velikih očekivanja pa sve razdoblja razočaranja i nezadovoljstva (Bostrom, 2014).

Kroz godine, John McCarthy je u široj javnosti stekao titulu oca umjetne inteligencije. On je ovaj vrlo kompleksan pojam pokušao pojednostaviti i pojasniti kao znanstvenu disciplinu koja se bavi izradom računalnih strojeva i programa čije se ponašanje može protumačiti kao inteligentno (McCarthy, 2004). Dakako, ponuđene su i kompleksnije i opširnije definicije poput one koju navodi Europski revizorski sud (2024) u službenom dokumentu pod nazivom Ambicije Europske unije u području umjetne inteligencije, a koji kaže kako se pod pojmom umjetne inteligencije stavljaju sustavi koji pokazuju inteligentno ponašanje tako što analiziraju

svoje okruženje i izvode radnje radi postizanja određenih ciljeva uz određeni stupanj autonomije. Nadalje, pojam umjetne inteligencije obuhvaća razne tehnologije koje se mijenjaju i u slučaju kojih se razvija sinergija s drugim tehnološkim trendovima (od robotike, gomile podataka, računalstva u oblaku, računalstva visoke razine učinkovitosti, fotonike i neuroznanosti). Osobito velik i revolucionaran pomak u razvoju umjetne inteligencije postignut je razvojem algoritama za strojno učenje koji uče iz velikih baza podataka i kroz vrijeme uspijevaju povećati svoju točnost (Seseri i sur., 2024).

Sada već gotovo zastarjela definicija koju Europska komisija navodi 2018. godine ističe kako sustavi umjetne inteligencije mogu biti softverski i djelovati u virtualnom okruženju, kao što su primjerice virtualni glasovni asistenti, biometrijski sustavi za prepoznavanje lica ili glasa, ali da mogu biti i hardverski sustavi kao što su primjerice roboti, autonomna vozila, dronovi i drugi uređaji koji imaju mogućnost biti spojeni na internet i komunicirati s drugim uređajima (Internet stvari / engl. *Internet of Things*) (Europska komisija, 2018).

U ožujku 2024. godine Europski parlament službeno je usvojio Akt o umjetnoj inteligenciji (Europski parlament, 2024a). Navedeni dokument predstavljao je važan pravni okvir za regulaciju proizvodnje, prihvatanja i korištenja umjetne inteligencije. Premda je riječ o aktu koji je na snazi u Europskoj uniji, on je globalno bitan jer se tehnološki divovi iz SAD-a ili Kine zbog njegove primjene moraju prilagođavati europskim korisnicima. U Aktu o umjetnoj inteligenciji stoji kako sustav umjetne inteligencije predstavlja sustav temeljen na stroju koji je dizajniran za rad s različitim razinama autonomije i koji može pokazivati prilagodljivosti nakon postavljanja te koji, za eksplicitne ili implicitne ciljeve, na temelju ulaznih podataka koje prima zaključuje kako generirati rezultate kao što su predviđanja, sadržaj, preporuke ili odluke koje mogu utjecati na fizička ili virtualna okruženja (*EU Artificial Intelligence Act*, 2024).

Umjetna inteligencija predstavlja važnu stratešku kariku za konkurentnost Europske unije na globalnom tržištu. To dokazuje i prijedlog Europske komisije iz 2021. godine pod nazivom „Put u digitalno desetljeće“ u kojem je vidljivo kako je cilj Europske unije do 2030. godine napraviti digitalnu transformaciju gospodarstvu u kojem će (Europska komisija, 2021):

- 75 % poduzeća u Europskoj uniji koristiti računalstvo u oblaku, umjetnu inteligenciju i velike podatke
- povećanjem financijskih sredstava za inovatore biti udvostručen broj „jednoroga“ u Europskoj uniji

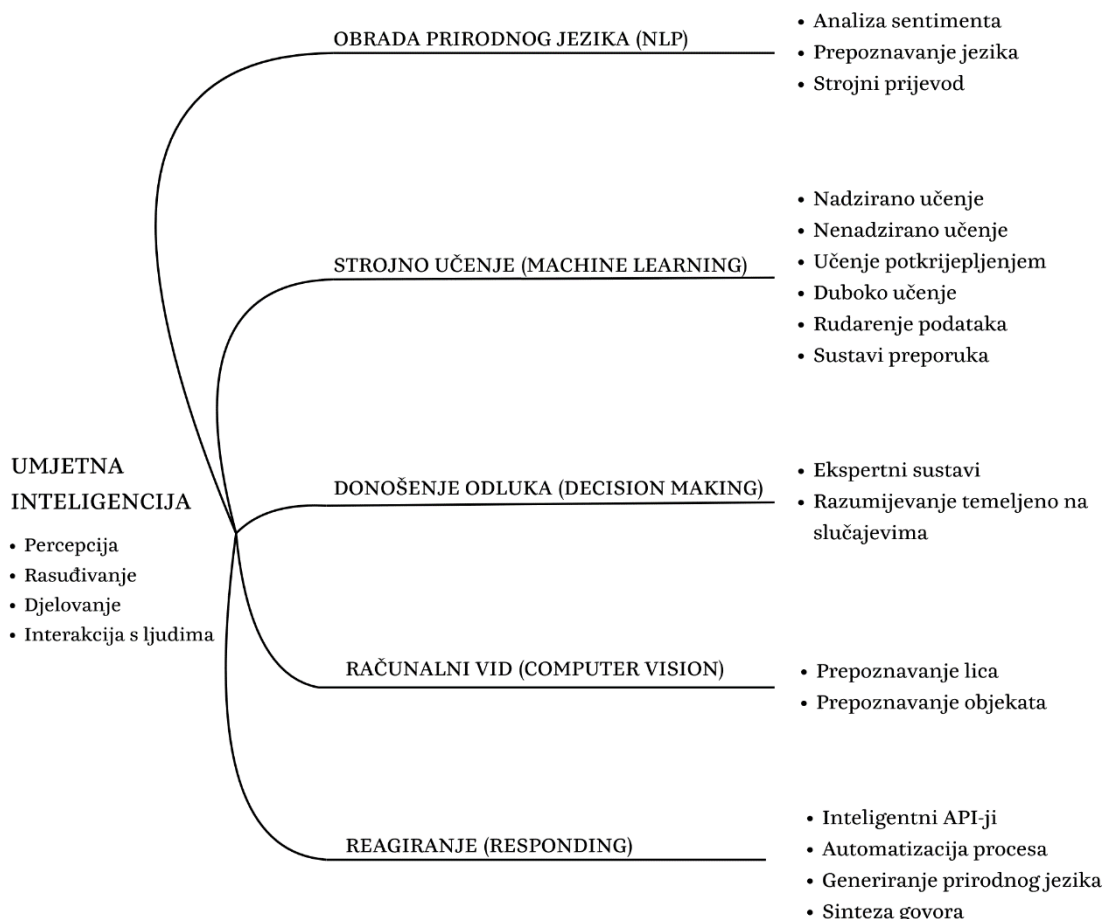
- više od 90 % malih i srednjih poduzeća doseći barem osnovnu razinu digitalnog intenziteta.

S druge strane, Republika Hrvatska u Strategiji digitalne Hrvatske za razdoblje do 2032. u značajno manjoj mjeri naglašava umjetnu inteligenciju kao tehnologiju oko koje će graditi svoju konkurentnost i digitalno društvo. Ipak, iako se u glavnim smjernicama za digitalnu tranziciju gospodarstva umjetna inteligencija ne navodi izričito, naglašena je u samom uvodu dokumenta: „Ova Strategija u sljedećem desetljeću pretpostavlja primjenu naprednih tehnologija kao što su 5G, 6G, umjetna inteligencija (engl. *artificial intelligence* – AI), strojno učenje (engl. *machine learning*), računalstvo u oblaku (engl. *cloud computing*), tehnologija velikih podataka (engl. *Big data*) i tehnologija lančanih blokova (engl. *blockchain*) u javnom i privatnom sektoru, ali i ostaje otvorena za implementaciju nekih budućih disruptivnih tehnologija koje će se pojaviti u promatranom periodu“ (Središnji državni ured za razvoj digitalnog društva, 2023).

U službenom dokumentu Bijele kuće pod nazivom „Izvršna uredba o sigurnom i pouzdanom razvoju i korištenju umjetne inteligencije“ (engl. *Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence*) pojam umjetna inteligencija ili AI definiran je u 15 USC 9401(3) kao sustav temeljen na stroju koji može, za određeni skup ciljeva definiranih od strane ljudi, davati predviđanja, preporuke ili odluke koje utječu na stvarno ili virtualno okruženje. Sustavi umjetne inteligencije koriste strojne i ljudske ulazne podatke kako bi percipirali stvarna ili virtualna okruženja, automatski apstrahiraju takve percepcije u modelu kroz analizu te koriste zaključivanje modela za formuliranje opcija za informacije ili djelovanje. Sustav umjetne inteligencije označava svaki podatkovni sustav, softver, hardver, aplikaciju, alat ili uslužni program koji u cijelosti ili djelomično radi koristeći umjetnu inteligenciju. (...) Izraz generativne umjetne inteligencije označava klasu modela umjetne inteligencije koji oponaša strukturu i karakteristike ulaznih podataka kako bi se generirao izvedeni sintetički sadržaj. To može uključivati slike, videozapise, zvuk, tekst i drugi digitalni sadržaj“ (The White House, 2023).

Kako je vidljivo na slici 2., ispod teksta, Mondal (2020) navodi glavna područja primjene te značajna polja umjetne inteligencije. Poseban naglasak stavlja na obradu prirodnog jezika (NLP – engl. *Natural language processing*), strojno učenje (ML – engl. *Machine learning*), donošenje odluka od strane umjetne inteligencije, računalni vid (engl. *Computer vision*) i reagiranje. Navedena slika daje jasan uvid u poimanje mogućnosti AI tehnologije u 2020. godini. Navedena se percepcija značajno mijenja u godinama koje su potom uslijedile, a gdje je snažan naglasak stavljen na mogućnost umjetne inteligencije da generira novi sadržaj. Premda je

obrada i generiranje prirodnog jezika komponenta koja je ekspertnim skupinama bila poznata i tada, šira javnost nije bila upoznata s tim mogućnostima umjetne inteligencije, osim kroz komunikaciju sa virtualnim asistentima poput Siri ili Alexe.



Slika 2. Primjene i polja umjetne inteligencije

Izvor: prilagođeno prema Mondalu (2020)

Kada se priča o modelima umjetne inteligencije pogonjenima dubokim neuronskim mrežama, tada se modeli mogu podijeliti u dvije skupine ovisno o različitim vrstama zadataka, a to su diskriminativni zadatci i generativni zadatci. Diskriminativni zadatci odlučuju o ulaznim podacima, kao što je prepoznavanje teksta, segmentiranje slika i slične klasifikacije, dok generativni zadatci uključuju stvaranje novih uzoraka podataka s obzirom na neke ulazne podatke. Generativni modeli prilagođeni su i osposobljeni za generiranje novih podataka, najčešće za prevođenje ili sažimanje tekstova, generiranje slika i zvukova ili odgovaranje na pitanja (WIPO, 2024).

2.3. Generativna umjetna inteligencija

Generativna umjetna inteligencija relativno je novo polje unutar umjetne inteligencije, a koje je snažno privukao pozornost javnosti u studenom 2022. godine kada je tvrtka OpenAI predstavila ChatGPT (Brühl, 2023). ChatGPT je najpopularniji alat generativne umjetne inteligencije (Statista, 2024), a kada mu se da zadatak da pojasni i definira što je to generativna umjetna inteligencija, rezultat generiranja je sljedeći (OpenAI, 2024):

- „Generativna umjetna inteligencija (*generative AI*) podskup je umjetne inteligencije koji se usredotočuje na stvaranje novih sadržaja – poput teksta, slika, glazbe pa čak i videozapisa koji nisu izravno kopirani iz podataka na kojima je model treniran. To se postiže korištenjem složenih modela, kao što su duboke neuronske mreže, koji mogu naučiti obrasce i strukture u podacima te ih koristiti za generiranje novih, originalnih izlaza.
- definicija: Generativna umjetna inteligencija odnosi se na vrstu umjetne inteligencije koja stvara nove podatke ili sadržaje imitirajući uzorke iz velikih skupova podataka na kojima je trenirana. Ovaj tip AI može autonomno generirati tekst, slike, zvukove ili druge oblike informacija koji su slični, ali ne identični, stvarnim podacima.
- primjeri generativne AI uključuju modele kao što su GPT (za generiranje teksta), DALL·E (za generiranje slika), i DeepFake (za stvaranje lažnih videozapisa)“.

Generativna umjetna inteligencija (akronim - Gen AI) koristi generativno modeliranje i napredak u dubokom učenju (engl. *Deep learning*) kako bi generirala raznovrstan sadržaj u velikom obujmu koristeći postojeće medije poput teksta, zvuka, videa, slika i drugih grafika (Jovanović & Campbell, 2022).

Pojednostavljeno, atribut generativna ispred pojma umjetna inteligencija označava proizvodnju ili stvaranje određenog sadržaja od strane određenih alata generativne umjetne inteligencije. Termin generativna umjetna inteligencija krovni je pojam za algoritme koji se mogu koristiti za stvaranje sadržaja, uključujući audio sadržaje, kodiranje, vizualne sadržaje, video sadržaje, simulacije te tekstualne sadržaje (Kalota, 2024).

Za razliku od klasičnih modela strojnog učenja koji prepoznavanjem uzoraka u podacima za obuku uče predviđati, klasificirati, davati personalizirane preporuke ili pružati podršku u donošenju odluka, generativna umjetna inteligencija omogućuje izrazito brzo stvaranje novih sadržaja na temelju korisničkog upita (Murugesan i Cherukuri, 2023). Generativna umjetna

inteligencija ima sposobnost učenja iz postojećih podataka na temelju kojih generira nove sadržaje. Rezultat generiranja često je vrlo realističan sadržaj koji zadržava karakteristike izvornog skupa podataka, ali ne ponavlja sadržaj (Gartner, 2023; Rois-Campos i sur., 2023).

Generativna umjetna inteligencija obuhvaća različite vrste od kojih je svaka prilagođena određenim zadacima, odnosno oblicima generiranja medija. Neke od najistaknutijih vrsta generativnih arhitektura AI modela su generativne protivničke mreže (GAN – engl. *Generative Adversarial Networks*), transformerski modeli (TRM – engl. *Transformer-based Models*), varijacijski autoenkoderi (VAE – engl. *Variational Autoencoders*) i difuzijski modeli (DM – engl. *Diffusion models*) (Sengar i sur., 2024).

Koncept generativnih protivničkih mreža (GAN) predstavljen je 2014. godine, a sastoji se od dvije neuronske mreže: generatora i diskriminatora (Goodfellow i sur., 2014). Navedene mreže zajedno sudjeluju u kompetitivnom procesu učenja pri čemu generator stvara sintetičke uzorke podataka, primjerice fotografije, videozapise ili glazbu, dok ih diskriminator procjenjuje kako bi razlikovao stvarne od lažnih primjera. Rezultati generativnih protivničkih mreža iznimno su realistični upravo zbog tog protivničkog učenja (Goyal i Mahmoud, 2024). Takav način učenja umjetne inteligencije rezultira i time da sustav kontinuirano poboljšava svoje sposobnosti generiranja podataka i tako proizvodi visokokvalitetne i realistične rezultate, zbog čega je sjajan za primjenu u svrhu generiranja i manipuliranja fotografijama, videozapisima, prepoznavanju i klasifikaciji objekata te obradi prirodnog jezika (Banh & Strobel, 2023; Aggarwal i sur., 2021). Generativne protivničke mreže pokazale su se izrazito uspješnim u zadacima kao što su prijenos stila (primjerice stil slike iz fotografije u crtež) ili povećanje podataka (kao primjerice stvaranje novih sintetičkih podataka za povećanje veličine i raznolikosti skupa podataka za obuku) (Stryker & Scapicchio, 2024).

Varijacijski autoenkoderi (VAE), još jedan od modela dubokog učenja, funkcioniraju na način da uče kako smanjiti ulazne podatke u manji prikaz, poznat i kao latentni prostor, koji se zatim koristi za rekonstrukciju izvornih podataka (Kingma & Welling, 2013; Goyal i Mahmoud, 2024). Optimiziranjem varijacijske donje granice vrijednosti podataka, varijacijski autoenkoderi mogu generirati nove uzorke koji nalikuju izvornoj distribuciji podataka. Tipičan primjer upotrebe varijacijskih autoenkodera može se vidjeti u sintetičkom generiranju i rekonstrukciji podataka kao što su slike, prepoznavanje anomalija i sustavima preporuka (Wei & Mahmood, 2021; Banh & Strobel, 2023). Stryker i Scapicchio (2024) navode kako su VAE modeli izrazito dobri pri analiziranju medicinskih slika zbog kvalitetnog otkrivanja anomalija.

Transformerski modeli (TM – engl. *Transformer-based models*) postali su temelj za mnoge napredne sustave koji se koriste za obradu prirodnog jezika i nasljedne modele. Riječ je o tipu neuronske mrežne arhitekture koji koriste mehanizam samopažnje i mehanizam višestruke pažnje za hvatanje dugoročnih ovisnosti i obrazaca u podacima, a zbog čega su idealni za obradu jezika na velikim bazama podataka (Bahn & Strobel, 2023; Vaswani i sur., 2017). Transformeri ili transformerski modeli često se koriste za izradu generativnih modela umjetne inteligencije kao što su generativni unaprijed obučeni transformeri (GPT – engl. *Generative Pre-trained Transformers*) koji su sposobni generirati kontekstualno relevantan tekst. Jedan od najpoznatijih alata generativne umjetne inteligencije, ChatGPT tvrtke OpenAI, temelji se na transformerskim modelima kako i njegovo samo ime sugerira (GPT – *Generative Pre-trained Transformer*) (Sengar i sur., 2024). Transformerski modeli srž su većine današnjih generativnih AI alata, posebice onih najpopularnijih kao što su ChatGPT, Copilot, BERT, Gemini. Transformerski modeli omogućuju brzu obuku te su izvrsni u obradi prirodnog jezika i razumijevanju prirodnog jezika zbog čega mogu generirati duže nizove podataka, odgovarati na pitanja, pisati pjesme, članke, kodirati, razumjeti i računati složenije tablice te još mnogo toga, a što je iznimno važno – to rade s visokom preciznošću (Stryker & Scapicchio, 2024).

Latentni difuzijski modeli (LDM) ili jednostavno difuzijski modeli temelje se na konceptima usklađivanja rezultata uklanjanja šuma i kontrastivnog divergiranja za učenje procesa generiranja stohastičkih podataka (Ho i sur., 2020; Rombach i sur., 2022). U svom istraživanju Dhariwal & Nichol (2021) navode kako je LDM nadmašio GAN u sintezi slika, a nakon čega se težište znanstvenih i stručnih istraživanja preusmjerilo na istraživanja difuzijskih modela što je dovelo do sjajnih rezultata u različitim zadacima generativnog modeliranja. Proces obuke difuzijskih modela specifičan je po tome što osigurava višu razinu stabilnosti nego li GAN-ov proces obuke, a njegova sposobnost generiranja širokog spektra visokokvalitetnih uzoraka bolja je od one VAE-ove (Liu i sur., 2023). Kako navode Stryker i Scapicchio (2024), difuzijskim je modelima potrebno više vremena za treniranje nego li VAE ili GAN modelima, ali u konačnici nude detaljniju kontrolu nad izlazom, posebno kada je riječ o alatima za generiranje visokokvalitetnih slika. DALL-E, alat koji je razvio OpenAI, primjer je alata kojega pokreće difuzijski model.

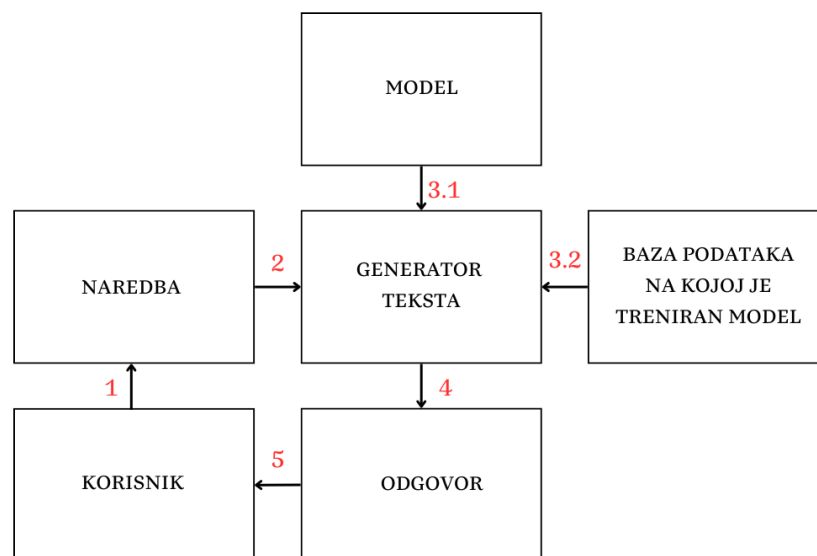
Iako se ovaj rad usredotočuje na potpuni potencijal generativne umjetne inteligencije, naglasak se mora staviti i na velike jezične modele kao najprepoznatljivijem tipu alata generativne umjetne inteligencije. Veliki jezički modeli (LLM – engl. *Large Language Model*) podskup su

generativne umjetne inteligencije specijalizirani za generiranje teksta, osposobljeni da generiraju tekst i odgovaraju na pitanja tekstualno (Corchado i sur., 2023). Mnoge tvrtke su tijekom 2010-ih implementirale različite jezične modele kako bi poboljšale svoje mogućnosti razumijevanja prirodnog jezika (NLU) i obrade prirodnog jezika (NLP), a to se događalo usporedno s napretkom u strojnom učenju, dubokom učenju, algoritmima, neuronskim mrežama i transformerskim modelima koji pružaju arhitekturu za velike jezične modele. Riječ je o modelima koji su obučavani na ogromnim bazama podataka i mogu se upotrebljavati za razne zadatke. Dizajnirani su da daju dojam razumijevanja teksta i generiraju tekst poput čovjeka, da imaju sposobnost zaključivanja iz konteksta, da pružaju relevantne odgovore, prevode ili sažimaju tekstove, odgovaraju na pitanja, pomažu u kreativnom pisanju, kodiranju te izvršavaju još mnoge druge zadatke. Javnosti je LLM dostupan kroz alate poput ChatGPT-a (IBM, 2024). Ipak, spomenuti ChatGPT i njemu slični (npr. Gemini, Claude, DeepSeek) više su od velikih jezičnih modela, riječ je o aplikacijama koje obrađuju značajno kompleksnije zadatke od generiranja teksta i razumijevanja konteksta. Riječ je o aplikacijama generativne umjetne inteligencije koji pružaju veći opseg funkcionalnosti nego što pojam velikih jezičnih modela predstavlja u teoriji (Wiemer, 2023). Iz navedenog razloga javlja se i novi pojam – multimodalni veliki jezični modeli (MLLM – engl. *multimodal large language model*). Riječ je o modelima koji nadilaze ograničenja isključivo tekstualnog unosa i mogu pristupiti znanju iz više modaliteta i tako mogu potpunije komunicirati (WIPO, 2024).

Corchado i sur. (2023) navode kako su veliki jezični modeli revolucionirali područje obrade prirodnog jezika, a karakteristični elementi velikih jezičnih modela su sljedeći:

- veliki broj parametara - veliki jezični modeli nose epitet „veliki“ zbog velikog broja parametara koji imaju, pa tako GPT-3 ima 175 milijardi parametara, a Googleov PaLM ima čak 540 milijardi parametara. Dokaz kako modeli postaju sve veći govori podatak podatak kako nasljednik, ChatGPT-4, ima 1,8 bilijuna parametara (Howarth, 2025).
- trenirani su na velikim bazama podataka: LLM-ovi su trenirani na ogromnim skupovima podataka koji obuhvaćaju veliki dio interneta, uključujući i knjige, mrežne stranice, znanstvene i stručne članke te još mnoge druge izvore znanja. To im omogućuje širok opseg znanja o jeziku, ali i raznim temama.
- sposobnost generiranja teksta - veliki jezični modeli mogu generirati tekst koji je koherentan i tečan te ga se često ne može razlikovati od teksta koji je napisao čovjek. Veliki jezični modeli mogu pisati eseje, pjesme, pisati znanstvene i stručne članke, poeziju i još mnoge druge forme teksta.

- transferno učenje – nakon što su istrenirani na velikim bazama podataka, veliki jezični modeli mogu se „fino doraditi“ (engl. *fine tuning*) za specifične zadatke uz relativno malu količinu podataka specifičnih za zadatak
- korištenje arhitekture transformera – većina modernih LLM-ova, poput GPT-a i BERT-a, temelji se na arhitekturi transformera koji koristi mehanizme pozornosti za hvatanje odnosa unutar podataka
- multimodalne mogućnosti – premda su veliki jezični modeli primarno usmjereni na tekst, noviji veliki jezični modeli imaju multimodalne mogućnosti, što znači da mogu razumjeti i generirati više vrsta podataka
- generalizacija kroz zadatke – bez potrebe za specifičnim promjenama u arhitekturi, veliki jezični modeli mogu obavljati razne zadatke, dovoljno je samo dati odgovarajući i dovoljno specificiran upit
- etički izazovi i pristranost – budući da se veliki jezični modeli dobrim dijelom uče iz javno dostupnih internetskih podataka, mogu usvojiti i pristranosti prisutne u tim podacima. Ova karakteristika velikih jezičnih modela jedna je od najkontroverznijih i izaziva brojne rasprave o etičnosti uporabe takvih sustava.



Slika 3. Pojednostavljeni prikaz funkcioniranja velikih jezičnih modela

Izvor: Kirova i sur. (2023)

Postoji više vrsta velikih jezičnih modela, Corchado i sur. (2023) navode četiri glavne vrste:

- 1) Autoregresivni modeli

- GPT (*Generative Pre-Trained Transformer*): razvijen od strane OpenAI-ja, GPT je autoregresivni model koji generira tekst riječ po riječ. GPT-1 predstavljen je u lipnju 2018. godine i sadržavao je 117 milijuna parametara. GPT-2 predstavljen je u veljači 2019. godine i sadržavao je 1,5 milijardi parametara. GPT-3 predstavljen je u lipnju 2020. godine, a sadržavao je 175 milijardi parametara. Navedeni, GPT-3 postao je javno dostupan široj javnosti u studenom 2022. godine.
- 2) Dvosmjerni modeli za klasifikaciju
- BERT (*Bidirectional encoder representation from transformers*): razvijen od strane Googlea, BERT je model koji je treniran u dva smjera, što znači da u obzir uzima kontekst s obje strane riječi u rečenici. Ovaj tip velikih jezičnih modela posebno je koristan za zadatke čitanja s razumijevanjem i klasifikaciju teksta.
- 3) Modeli sekvenca-u-sekvencu
- T5 (*Text-to-text transfer transformer*): također razvijen od strane Googlea, T5 interpretira sve zadatke obrade jezika kao problem pretvaranja teksta u tekst. Zadatci poput prijevoda, sažimanja, odgovaranja na pitanja rješavaju se kao transformacije od ulaznog do izlaznog teksta.
 - BART (*Bidirectional and auto-regressive transformers*): razvijen od strane Facebook AI-ja. BART kombinira značajke BERT-a i GPT-a za zadatke generiranja i razumijevanja teksta.
- 4) Multimodalni modeli:
- CLIP (*Contrastive Language-Image Pre-training*) i Dall-E: oba modela razvijena su od strane OpenAI-ja i kombiniraju računalni vid i obradu prirodnog jezika. Dok CLIP razumije slike u kontekstu prirodnog jezika, Dall-E generira slike iz tekstualnog zapisa.
 - WU DAO: dubinski jezični model koji je razvijen od strane Pekinške akademije za umjetnu inteligenciju. Ima multimodalne značajke i treniran je na tekstualnim i slikovnim podacima, što mu omogućava rješavanje oba tipa zadataka. Ima ogroman broj parametara (1,75 bilijuna).

2.4. Primjena generativne umjetne inteligencije

Dolazak generativne umjetne inteligencije najavljuje velike prilike za primjenu u raznim industrijama. Organizacije istražuju kako se generativna umjetna inteligencija može koristiti za unapređenje poslovnih procesa, povećanje učinkovitosti i produktivnosti te otvaranje mogućnosti za potpuno nove proizvode, usluge i poslovne modele (Deloitte, 2024). Kako

navodi Microsoft u svojoj studiji (Taylor, 2024), generativna umjetna inteligencija postala je pokretač poslovnih rezultata u gotovo svim industrijama. Premda je relativno nova tehnologija, u 2024. godini čak 75 % zaposlenih ljudi koristi neku vrstu generativne umjetne inteligencije, dok je ta brojka u istraživanju iz 2023. godine bila 55 %. Nadalje, navedeno izvješće navodi kako je povrat investicije uložen u generativnu umjetnu inteligenciju čak 370 % (Taylor, 2024). S druge strane, Liu i Wang (2024) navode kako su ove brojke najsnažnije na azijskim tržištima, posebno u Indiji, Kini, Tajlandu i Indoneziji, gdje više od 90 % ispitanika navodi da koristi generativnu umjetnu inteligenciju u poslovne svrhe. Nadalje, Liu i Wang (2024) navode kako otprilike 12,5 % radnog stanovništva koristi ChatGPT u poslovne svrhe.

Istraživanje koje je proveo Indeed govori kako stadij razvijenosti generativne umjetne inteligencije u ovom trenutku nije dovoljno snažan da bi predstavljao ozbiljniju ugrozu ljudskoj radnoj snazi (Hering & Rojas, 2024). Generativna umjetna inteligencija općenito je snažna u pružanju teorijskog znanja o vještinama, ali u korištenju vještina za rješavanje problema nije toliko vješta. Posebno kod zanimanja koja zahtijevaju značajnu praktičnu upotrebu, primjerice upravljanje zrakoplovom ili kuhanje, korisnost generativne umjetne inteligencije je vrlo ograničena. Unatoč vrlo brzom razvoju navedene tehnologije, masovno zamjenjivanje ljudskih radnika u ovom trenutku nije realno. Ipak, određena zanimanja ugroženija su od drugih, a među njima se prvenstveno ističu radnici u računovodstvu, marketingu, razvoju softvera, administrativnoj podršci u zdravstvu te djelatnici i procjenitelji u osiguravajućim kućama (Hering & Rojas, 2024).

FlexOs (2024a) navodi kako je ChatGPT najprepoznatljiviji alat generativne umjetne inteligencije, čak 75 % ljudi čulo je za njega. Osim njega, vrlo prepoznatljivi alati generativne umjetne inteligencije koje prepoznaje opća populacija su Grammarly (35 %), Google Bard / Google Gemini (33 %), Microsoft Copilot (25 %), Adobe Firefly (23 %), Copy.AI (15 %), Dall-E (12 %), CodaAI (9 %), Notion AI (9 %), Jasper (8 %) i Midjourney (6,5 %). Isto istraživanje navodi kako je i poredak po pitanju aktivnog korištenja vrlo sličan, najzastupljeniji je ChatGPT (60 %), a prate ga Google Bard/Gemini (27 %), Grammarly (19 %), Microsoft Copilot (18 %), Adobe Firefly (12 %), Copy.AI (12 %), Dall-E (4 %), Midjourney (4 %), Notion AI (4 %), Coda AI (4 %) te Jasper (3 %). Kada su u pitanju alati generativne umjetne inteligencije, Statista (2024) također navodi kako je uvjerljivo najkorišteniji alat ChatGPT, dok je kod obrade vizualnih sadržaja uvjerljivo najkorišteniji alat Midjourney.

FlexOS (2024b) u istraživanju o najkorištenijim alatima generativne umjetne inteligencije navodi kako GPT-evi (LLM-ovi) čine 66 % korištenja takvih vrsta alata. Primjeri alata unutar

kategorije GPT-eva su ChatGPT, Bing AI, Bard/Gemini, Claude, Microsoft Copilot, a od čega ChatGPT čini čak 76 % unutar kategorije. Druga najznačajnija kategorija su alati za pisanje i uređivanje tekstova koji čine 9 % od ukupnog postotka korištenja alata generativne umjetne inteligencije, a primjeri takvih alata su Grammarly, Simplified AI, ZeroGPT, Copy.AI, WriteSonic, Undetectable AI, WordTune, Jasper, Quillbot, a od čega Grammarly čini 71 % unutar kategorije.

Treća najznačajnija kategorija sa 6 % unutar ukupnog korištenja alata generativne umjetne inteligencije su alati za obrazovanje kao što su Brainly, CourseHero, TurnItIn, ELSA, MagicSchool i Caktus, od čega je najzastupljeniji Brainly sa 42 % unutar kategorije. Četvrta najzastupljenija kategorija (5 %) su alati za kreiranje virtualnih persona i razgovor s njima kao što su Character.AI, AI Chatting, Chai, Candy AI, Anima, Romantic AI i drugi, među kojima je najkorišteniji upravo Character.AI koji čini 89 % unutar kategorije. Preostale značajne kategorije korištenja alata generativne umjetne inteligencije su alati za generiranje slika kao što su DeepAI, Midjourney, Canva, Looka, Adobe Firefly, Hotspot, NightCafe, ali i alati za istraživanja kao što su Perplexity, You i ChatPDF (FlexOS, 2024b).

Deloitte (2024) u svom istraživanju navodi kako su emocije koje ljudi najviše osjećaju pri korištenju generativne umjetne inteligencije uzbuđenje (62 %), fascinacija tehnologijom (46 %), neizvjesnost (30 %), povjerenje (17 %), iznenađenje (16 %), anksioznost (10 %), zbunjenost (8 %), strah (6 %), iscrpljenost (4 %) te ljutnja (1 %) što potvrđuje da je veći dio značajnijih emocija uključen u konceptualni model ovoga istraživanja.

FlexOS (2024a) u istraživanju o korištenju generativne umjetne inteligencije navodi kako čak 43,5 % ispitanika nikada ili gotovo nikada ne koriste alate generativne umjetne inteligencije. Navedeni podatci približni su onima koje spominju Biloš i Budimir (2024), koji navode kako je 59,9 % ispitanika barem probalo koristiti ChatGPT, dok McKinsey navodi brojku od 65 % (Singla i sur., 2024).

Podatci koje je izvješću Svjetske banke iznose Liu i Wang (2024) potvrđuju prethodne studije koje tvrde da je ChatGPT uvjerljivo najkorišteniji alat generativne umjetne inteligencije, ali i pokazuju kako su veliki jezični modeli općenito najkorišteniji vid upotrebe generativne umjetne inteligencije.

Tablica 1. Najkorišteniji alati generativne umjetne inteligencije

	Alat generativne umjetne inteligencije	Tip alata	Promet u ožujku 2024. (u milijunima)
1.	ChatGPT	Chatbot	2343,2
2.	Gemini	Chatbot	132,9
3.	Poe	Chatbot	43,4
4.	Perplexity	Chatbot	40,2
5.	Claude	Chatbot	32,3
6.	DeepAI	Chatbot	31,1
7.	Copilot	Chatbot	26,2
8.	Midjourney	Slika	24,7
9.	Prezi	Slika	18,0
10.	Nightcafe	Slika	13,9
11.	Leonardo	Slika	13,6
12.	Gamma	Slika	11,6
13.	Pixai	Slika	9,6
14.	Runway	Video	9,0
15.	Ideogram	Slika	8,9
16.	Playground	Slika	8,8
17.	Youchat	Chatbot	8,7
18.	Blackbox AI	Chatbot	8,7
19.	ChatPDF	Chatbot	6,8
20.	MaxAI	Chatbot	6,1

Izvor: Liu & Wang, 2024

Liu i Wang (2024) ističu kako postoji jaz među spolovima kada je u pitanju korištenje alata generativne umjetne inteligencije, posebno kada se uzmu u obzir pokazatelji korištenja nekih drugih internetskih alata. Ako se pogleda struktura korištenja alata poput tražilice Google gdje žene čine 48 % korisnika u ukupnoj strukturi ili Wikipedije gdje žene čine 52 % korisnika u ukupnoj strukturi, dolazi se do nekakvih očekivanih brojki jer žene čine 48 % ukupnih internetskih korisnika. Kod alata generativne umjetne inteligencije, zanimljivo je uočiti kako žene čine samo 33 % od ukupnog broja korisnika ChatGPT-a, ili pak samo 21 % korisnika Runwaya, najpopularnijeg alata generativne umjetne inteligencije za generiranje video

sadržaja. Ipak, izuzetak je Midjourney, alat generativne umjetne inteligencije za generiranje slikovnog sadržaja, gdje je udio ženskih korisnika 52 %.

Gotovo polovica korisnika *chatbotova* i velikih jezičnih modela su visokoobrazovani. Nadalje, alati generativne umjetne inteligencije privlače pretežno mlađe korisnike, ali *chatbotovi* i veliki jezični modeli već su široko prihvaćeni i među starijim korisnicima te im je dobna struktura slična Googleovoj. Koliko se ChatGPT koristi kao alat za poboljšanje radne produktivnosti govori i činjenica da se njegova upotreba tijekom vikenda smanji za 40 % u usporedbi s radnim danima, a ako se promatra upotreba ChatGPT-a na računalima, onda promet tijekom vikenda pada i do 50 % (Liu & Wang, 2024).

Prema informacijama koje je objavio američki Nacionalni institut za ekonomska istraživanja (Bick i sur., 2024), muškarci koriste generativnu umjetnu inteligenciju više nego žene. Isto istraživanje navodi kako alate generativne umjetne inteligencije u većoj mjeri koriste mlađe generacije i obrazovanije stanovništvo. Značajne razlike kod korištenja generativne umjetne inteligencije ističe i Deloitte (2024), premda naglašava kako se taj jaz smanjuje. Kada se priča o generacijskom jazu, Salesforce (2024) navodi kako 65 % korisnika alata generativne umjetne inteligencije čine milenijalci i pripadnici generacije Z, a čak 68 % nekorisnika takve tehnologije su pripadnici generacije X i *baby-boomeri*. Pew Research Center također navodi kako alate generativne umjetne inteligencije koriste češće mladi nego stariji, muškarci češće nego žene, obrazovaniji češće nego manje obrazovani, ali i nadodaju kako korisnici takve tehnologije žive u gradovima i prigradskim naseljima, a nešto rjeđe u ruralnim sredinama (Lin & Parker, 2025). Banka federalnih rezervi New Yorka u svom velikom istraživanju potvrđuje ranije navedene tvrdnje te ističe kako postoje statistički značajne razlike između muškaraca i žena kod korištenja generativne umjetne inteligencije, ali i nadodaju da takvu tehnologiju češće koriste mladi, visokoobrazovani i oni s većim primanjima (Aldaroso i sur., 2024a; Aldaroso i sur., 2024b). Podatci Microsoftovog istraživanja provedenog 2023. godine također potvrđuju kako među korisnicima generativne umjetne inteligencije ima više muškaraca, više mladih i više visokoobrazovanih (Microsoft, 2024).

Nadalje, istraživanje Nacionalnog instituta za ekonomska istraživanja (Bick i sur., 2024) navodi kako je gotovo 40 % odraslih barem probalo koristiti alate generativne umjetne inteligencije, a od tih korisnika 32 % njih da ju koristi barem jednom tjedno i 10 % njih da ih koristi svakodnevno. Isto istraživanje navodi kako alate generativne umjetne inteligencije češće koriste visokoobrazovani i oni čija su zanimanja tradicionalno više vezana za rad s tehnologijom (pisanje, administracija, prevođenje), ali da čak 22 % radnika u klasi plavih ovratnika koristi

alate generativne umjetne inteligencije. Alati koji se najviše koriste su ChatGPT (28,1 %), Gemini (16,1 %) i Microsoft Copilot (14,1 %).

Bick i sur. (2024) navode kako se generativna umjetna inteligencija usvaja brže nego prethodne tehnologije. Premda je usvajanje generativne umjetne inteligencije slično usvajanju osobnih računala tijekom 1980-ih, razlika se vidi kod mlađe i obrazovanije populacije koja sada brže usvaja generativnu umjetnu inteligenciju. Nadalje, Bick i sur. (2014) navode kako više od 20 % ispitanika ističe da im generativna umjetna inteligencija uštedi više od četiri sata tjedno, a trećina ispitanika navodi kako im alati generativne umjetne inteligencije uštede do sat vremena tjedno. Ipak, trećina ispitanika koji svakodnevno koriste generativnu umjetnu inteligenciju prijavljuju uštedu vremena veću od četiri sata tjedno. Kada se priča o intenzitetu korištenja, gotovo polovica korisnika navodi da alate generativne umjetne inteligencije koristi između 15 i 59 minuta dnevno, a 32 % ispitanika navodi da koristi više od 60 minuta dnevno. U pravilu, istraživanje pokazuje kako ispitanici koji više puta u tjednu koriste alate generativne umjetne inteligencije prijavljuju intenzivnije korištenje tijekom dana, ali i veću uštedu vremena zbog samog korištenja navedenih alata. To ukazuje na to da bi veća primjena alata generativne umjetne inteligencije mogla imati ključnu ulogu u transformaciji radnih procesa i optimizaciji radnog vremena.

Kopal i sur. (2024) navode kako se alati generativne umjetne inteligencije najviše upotrebljavaju za pisanje teksta (58 %), prijevod teksta (55 %), čavljanje (53 %), pronalazak teško dostupnih informacija (53 %), generiranje slika i videa (41 %), generiranje ideja (36 %), pisanje eseja/seminara/radova (27 %), pisanje računalnog koda (22 %). Globalno istraživanje koje je provedeno gotovo godinu dana ranije (Microsoft, 2024) također navodi kako se alati generativne umjetne inteligencije najviše koriste za prijevod teksta (43 %), pronalazak informacija (36 %), kao pomoć za pisanje studentskih radova (31 %), zabavu i čavljanje (29 %) te generiranje slika i videa (27 %).

Harvard Business Review navodi kako pojedinci sve više koriste alate generativne umjetne inteligencije kao terapeuta, asistenta koji im pomože u organizaciji privatnog i poslovnog života, ali koriste ih i za učenje, edukaciju, zabavu te kao savjetnika za zdraviji život. S obzirom na to da alati za izradu vizualnih sadržaja u 2025. godini bilježe velikih skok u kvaliteti generiranog sadržaja, tako bilježe i rast korištenja pa takve alate pojedinci koriste da generiraju ideje, izrade ili urede fotografije te takvim sadržajem zabave svoju djecu. Ukratko, u 2025. godini postoji trend korištenja generativne umjetne inteligencije za široki spektar aktivnosti koji

poboljšava kvalitetu privatnog života pojedinca, što u prethodnim godinama nije bio slučaj jer se spomenuta tehnologija dominantno koristila za poslovne zadatke (Zao-Sanders, 2025).

2.5. Etička perspektiva uporabe generativne umjetne inteligencije

Jedno od najvažnijih pitanja kod upotrebe umjetne inteligencije svakako je etičko pitanje. Etika umjetne inteligencije bavi se moralnim načelima koja upravljaju praktičnom primjenom umjetne inteligencije. Etika uključuje pronalaženje osjetljive ravnoteže između tehnološkog napretka i očuvanja ljudskih vrijednosti poput privatnosti, pravednosti, transparentnosti i odgovornosti. Iznimno je važno osigurati da odluke u kojima sudjeluju i sustavi umjetne inteligencije ne odražavaju ili pojačavaju društvene predrasude, a posebno je to bitno u područjima poput pravosuđa, zdravstva, socijalnih politika ili odluka o zapošljavanju (Kirova i sur. 2023).

Hagendorff (2024) navodi kako su teme koje najviše zabrinjavaju korisnike umjetne inteligencije, u etičkom smislu, pravednost i pristranost sustava umjetne inteligencije, sigurnost, štetni i opasni sadržaji koje takvi sustavi generiraju, halucinacije i neistine, privatnost, kibernetički kriminal, gubitak radnih mjesta, netransparentnost funkcioniranja sustava umjetne inteligencije te utjecaj generativne umjetne inteligencije na umjetnost, autorska prava, obrazovanje, pisanje, evaluacije i druga srodna područja. Na tom tragu su i rezultati koje navode Al-kfairy i sur. (2024), a riječ je zabrinutosti po pitanju autorstva i akademskog integriteta, regulatornih i pravnih pitanja, privatnosti, povjerenja i pristranosti, dezinformacija i *deepfake* sadržaja, transparentnosti i odgovornosti, autentičnosti te društvenih i ekonomskih utjecaja. Microsoft (2024) također navodi kako su strahovi javnosti najviše vezani za generiranje internetskih prevara, *deepfake* sadržaja, seksualno i online uznemiravanje, halucinacije sustava, privatnost podataka, pristranost sustava umjetne inteligencije, ali i za stvaranje previše bliskog odnosa sa takvim naprednim sustavima.

U dokumentu koji je objavila Europska komisija (2019), pod nazivom „Etičke smjernice za pouzdanu umjetnu inteligenciju“, pouzdana umjetna inteligencija ima tri sastavnice koje trebaju biti ispunjene tijekom cijelog životnog ciklusa sustava (Europska komisija, 2019):

- „(a) trebala bi biti zakonita i poštovati sve primjenjive zakone i propise
- (b) trebala bi biti etična i osigurati poštovanje etičkih načela i vrijednosti
- (c) trebala bi biti otporna i iz tehničke i iz socijalne perspektive jer sustavi umjetne inteligencije, čak i s dobrim namjerama, mogu uzrokovati nenamjernu štetu“.

Ključne smjernice za prvu točku navode kako je nužno razvijati i koristiti sustave umjetne inteligencije na način kojim se poštuju načela poštovanja ljudske autonomije, sprečavanja nastanka štete, pravednosti i objašnjivosti postupaka. Nužno je posebnu pozornost obratiti na situacije koje uključuju ranjive skupine kao što su djeca ili osobe s invaliditetom te druge skupine u nepovoljnom položaju i u riziku od društvene isključenosti. Premda sustavi umjetne inteligencije društvu i pojedincima pružaju znatne koristi, oni sa sobom nose određene rizike te mogu imati negativne učinke. Takve negativne učinke ponekad je teško predvidjeti te zbog toga tvorci sustava umjetne inteligencije moraju biti fleksibilni i spremni po potrebi donijeti odgovarajuće mjere za smanjenje tih štetnih posljedica (Europska komisija, 2019).

Ključne smjernice za drugu točku navode kako je nužno osigurati da razvoj, uvođenje i upotreba sustava umjetne inteligencije ispunjava zahtjeve za pouzdanu umjetnu inteligenciju, kao što su ljudsko djelovanje i nadzor, tehnička otpornost i sigurnost, privatnost i upravljanje podacima, transparentnost, raznolikost, nediskriminacija i pravednost, dobrobit za okoliš i društvo te odgovornost. Nužno je da se razmotre sve tehničke i netehničke metode kako bi se osigurala provedba tih zahtjeva te kako bi se sve dionike na jasan i proaktivan način informiralo o sposobnostima i ograničenjima sustava umjetne inteligencije. Ključno je olakšati sljedivost i provjerljivost sustava, posebno u kritičnim uvjetima ili situacijama (Europska komisija, 2019).

Nadalje, ovaj dokument naglašava važnost korištenja procjene pouzdanosti umjetne inteligencije tijekom cijelog životnog ciklusa sustava umjetne inteligencije, od razvoja i implementacije do same upotrebe. Ključno je tu procjenu prilagoditi specifičnoj namjeni sustava s obzirom na to da univerzalni popis za procjenu ne postoji. Osiguravanje pouzdane umjetne inteligencije ne podrazumijeva samo formalno ispunjavanje zahtjeva, već podrazumijeva stalan proces prepoznavanja i vrednovanja rješenja umjetne inteligencije. Cilj je pritom postići bolje rezultate za sve uključene strane. (Europska komisija, 2019).

3. MODELI PRIHVAĆANJA TEHNOLOGIJE

3.1. UTAUT

3.1.1. Razvoj modela

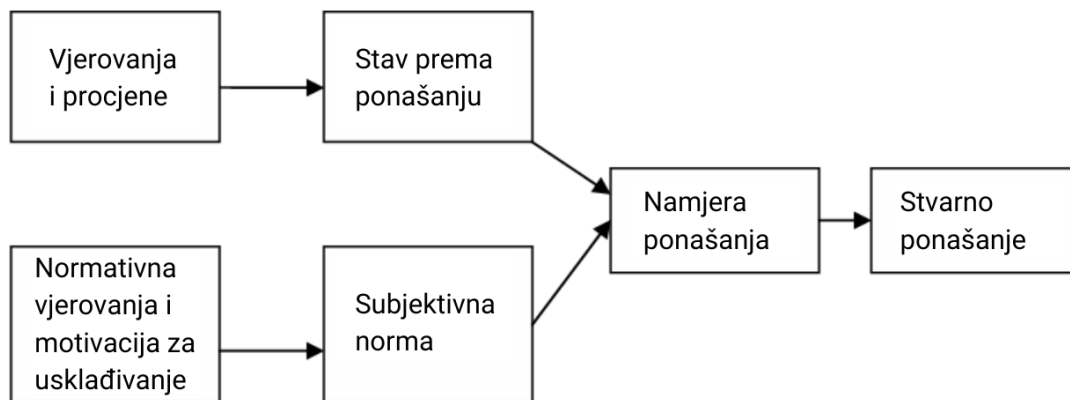
Akronim UTAUT2 dolazi od termina *Unified Theory of Acceptance and Use of Technology 2* što se na hrvatskom jeziku prevodi kao Ujedinjena teorija prihvaćanja i korištenja tehnologije 2 (Venkatesh i sur., 2012). Prva verzija modela objavljena je 2003. godine (Venkatesh i sur., 2003), a druga verzija modela objavljena je 2012. godine (Venkatesh i sur., 2012). Model je razvijao profesor Viswanath Venkatesh sa Sveučilišta Virginia Tech. Svjetski priznati znanstvenik koji je s više od 180 000 citiranja na Google znalcu i s više od 60 000 citiranja u Web of Science bazi jedan od 500 najcitiranijih znanstvenika na svijetu te najcitiraniji znanstvenik na svijetu u području informacijskih sustava (Venkatesh, 2024). Neki od znanstvenih projekata po kojima je profesor Venkatesh prepoznatljiv su razvoj modela TAM2 u radu „*A theoretical extension of the technology acceptance model: Four longitudinal field studies*“ te razvoj modela TAM3 u radu „*Technology acceptance model 3 and a research agenda on interventions*“. Osim spomenutih TAM2 i TAM3 modela, profesor Venkatesh prepoznat je po modelu UTAUT koji je predstavio 2003. godine u radu „*User acceptance of information technology: Toward a unified view*“ te unaprijeđenog i proširenog modela pod nazivamo UTAUT2 objavljenog 2012. godine u radu „*Consumer acceptance and use of information technology: extending the unified theory of acceptance and use of technology*“.

O izvornom radu o prvoj jedinstvenoj teoriji prihvaćanja i korištenja tehnologije (u nastavku rada UTAUT) Venkatesh navodi da jedinstvena teorija zapravo podrazumijeva ujedinjavanje nekoliko uspješnih teorija i modela koji su prethodili razvoju UTAUT modela, a to su (Venkatesh i sur., 2003):

- TRA – *engl. Theory of Reasoned Action* / hrv. Teorija razumnog djelovanja
- TAM – *engl. Technology Acceptance Model* / hrv. Model prihvaćanja tehnologije
- MM – *engl. Motivational Model* / hrv. Motivacijski model
- TPB – *engl. Theory of Planned Behavior* / hrv. Teorija planiranog ponašanja
- C-TAM-TPB – *engl. Combined TAM and TPB* / hrv. Kombinirani TAM i TPB
- MPCU – *engl. Model of PC Utilization* / hrv. Model korištenja osobnih računala
- IDT – *engl. Innovation Diffusion Theory* / hrv. Teorija difuzije inovacija
- SCT – *engl. Social Cognitive Theory* / hrv. Socijalno-kognitivna teorija.

Navedeni modeli međusobno su se nadopunjavali i sustavno usavršavali te sve bolje objašnjavali faktore koji utječu na prihvaćanje novih tehnologija što je bilo vrlo važno u doba tehnološke revolucije sa kraja 20. stoljeća kada je većina spomenutih modela i nastajala. Venkatesh i sur. (2003) empirijski su usporedili navedenih osam modela te na temelju konceptualnih i empirijskih sličnosti među modelima formulirali Ujedinjenu teoriju prihvaćanja i korištenja tehnologije.

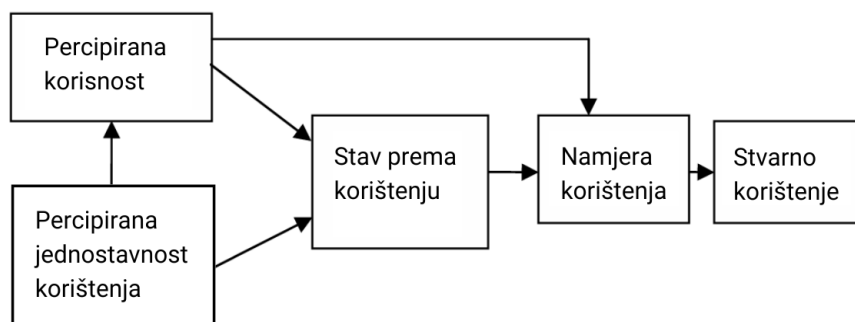
Teorija razumnog djelovanja (TRA) počiva na socijalnoj psihologiji i jedna je od osnovnih i najutjecajnijih teorija ljudskog ponašanja, a koristi se za predviđanje širokog spektra ponašanja (Fishbein & Ajzen, 1975; Venkatesh i sur., 2003; Sheppard i sur., 1988). Teorija razumnog djelovanja temelji se na pretpostavci da su pojedinci racionalni donositelji odluka koji neprestano vrednuju svoja vjerovanja o ponašanju u procesu oblikovanja svog stava prema tom ponašanju (Fishbein & Ajzen, 1975; Li, 2010). Davis i sur. (1989) primijenili su TRA na individualno prihvaćanje tehnologije i otkrili da je objašnjenje varijance modela u kontekstu tehnologije bila slična kao u kontekstu drugih ponašanja. Teorija razumnog djelovanja korištena je u raznim drugim kontekstima kao što je prihvaćanje tehnologije u poljoprivredi, marketingu, medicini, informatičkom poslovanju (Dissanayake, 2022; Xhang & Guo, 2013; Liker & Sindi, 1997; Fitzmaurice, 2005). Temeljni konstrukti TRA modela su Stav prema ponašanju (engl. *Attitude toward behavior*) i Subjektivna norma (engl. *Subjective norm*). Teorija razumnog djelovanja stavlja naglasak na stavove, uvjerenja i društvene norme zbog čega je pridonio razvoju različitih modela kao što su TAM i UTAUT te UTAUT2 (Venkatesh & Bala, 2008; Venkatesh i sur., 2012; Zhang i sur., 2014). Važan element ovog modela koji je kasnije preuzet od njegovih nasljednika u smislu proučavanja prihvaćanja tehnologije jest konstrukt namjere ponašanja i stvarnog ponašanja. Fishbein i Ajzen (1975) navode kako namjera pojedinca u ponašanju određuje njegovo ili njezino stvarno ponašanje. Odnos je to koji su kasnije preuzeli razni modeli, a među njima i oni koji istražuju prihvaćanje i korištenje informacijskih tehnologija.



Slika 4. Prikaz modela – Teorija razumnog djelovanja

Izvor: prilagođeno prema Fishbeinu & Ajzenu (1975)

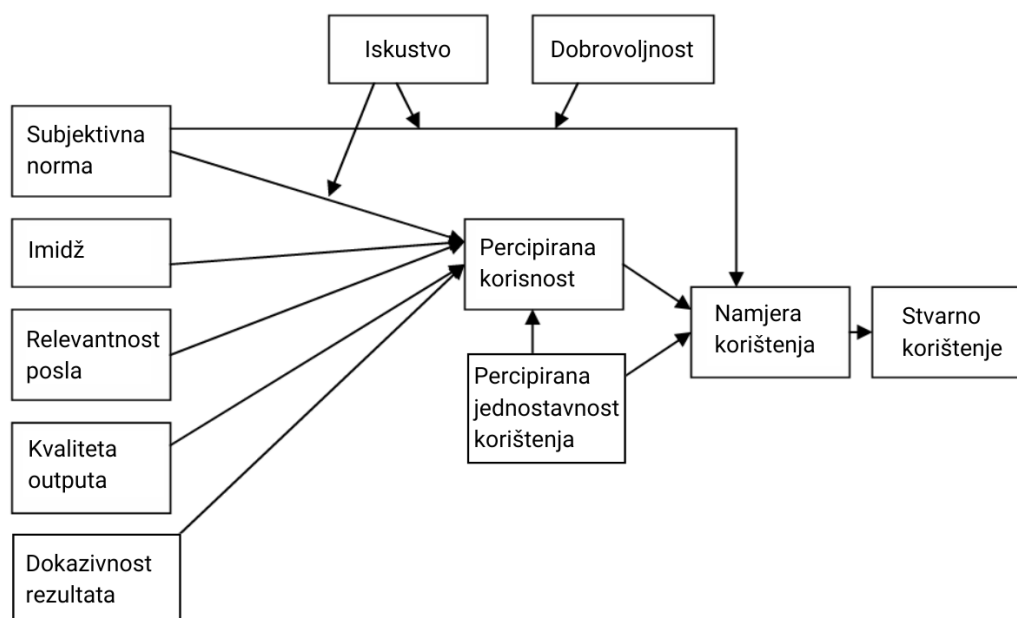
Model prihvatanja tehnologije (TAM) je model prilagođen kontekstu informacijskih sustava i dizajniran je za predviđanje prihvatanja i korištenja informacijske tehnologije u poslovnom okruženju, kroz godine je korišten u mnogobrojnim kontekstima i na različitim skupinama korisnika. Za razliku od teorije razumnog djelovanja, u modelu prihvatanja tehnologije isključen je konstrukt stavova koji je zamijenjen konstruktima percipirane jednostavnosti korištenja i percipirane korisnosti na namjeru pojedinca da koristi tehnologiju, a pri čemu je izbačena subjektivna norma kao medijator (Davis, 1989). Konstrukte percipirane korisnosti i percipirane jednostavnosti korištenja Davis (1989) vidi kao dva najvažnija individualna čimbenika kod korištenja informacijske tehnologije, a pri čemu je percipirana korisnost najsnažniji prediktor namjere pojedinca da koristi određenu informacijsku tehnologiju (Li, 2010).



Slika 5. Model prihvatanja tehnologije (TAM)

Izvor: prilagođeno prema Davisu (1989)

Nadograđena verzija TAM-a pod nazivom TAM2 proširuje izvorni model s konstruktima socijalnih utjecaja kao što su konstrukti subjektivne norme i imidža, ali i moderatorom dobrovoljnosti. Nadalje, specifičnost TAM2 modela je proširenje s konstruktima relevantnosti posla, kvalitete *outputa*, dokazivosti rezultata, ali i ranije spomenutim imidžom te subjektivnom normom kao prediktorima percipirane korisnosti, odnosno neizravnim prediktorima namjere korištenja i stvarnog korištenja tehnologije (Venkatesh & Davis, 2000). U TAM2 dodan je konstrukt subjektivne norme jer autori vjeruju kako ona ima izravan utjecaj na namjeru ponašanja kroz mehanizam usklađivanja. Ukoliko pojedinac vjeruje da važan društveni akter ima sposobnost kazniti određeno nepoželjno ponašanje ili nagraditi poželjno, doći će do učinka usklađivanja (Venkatesh & Davis, 2000; Li, 2010).

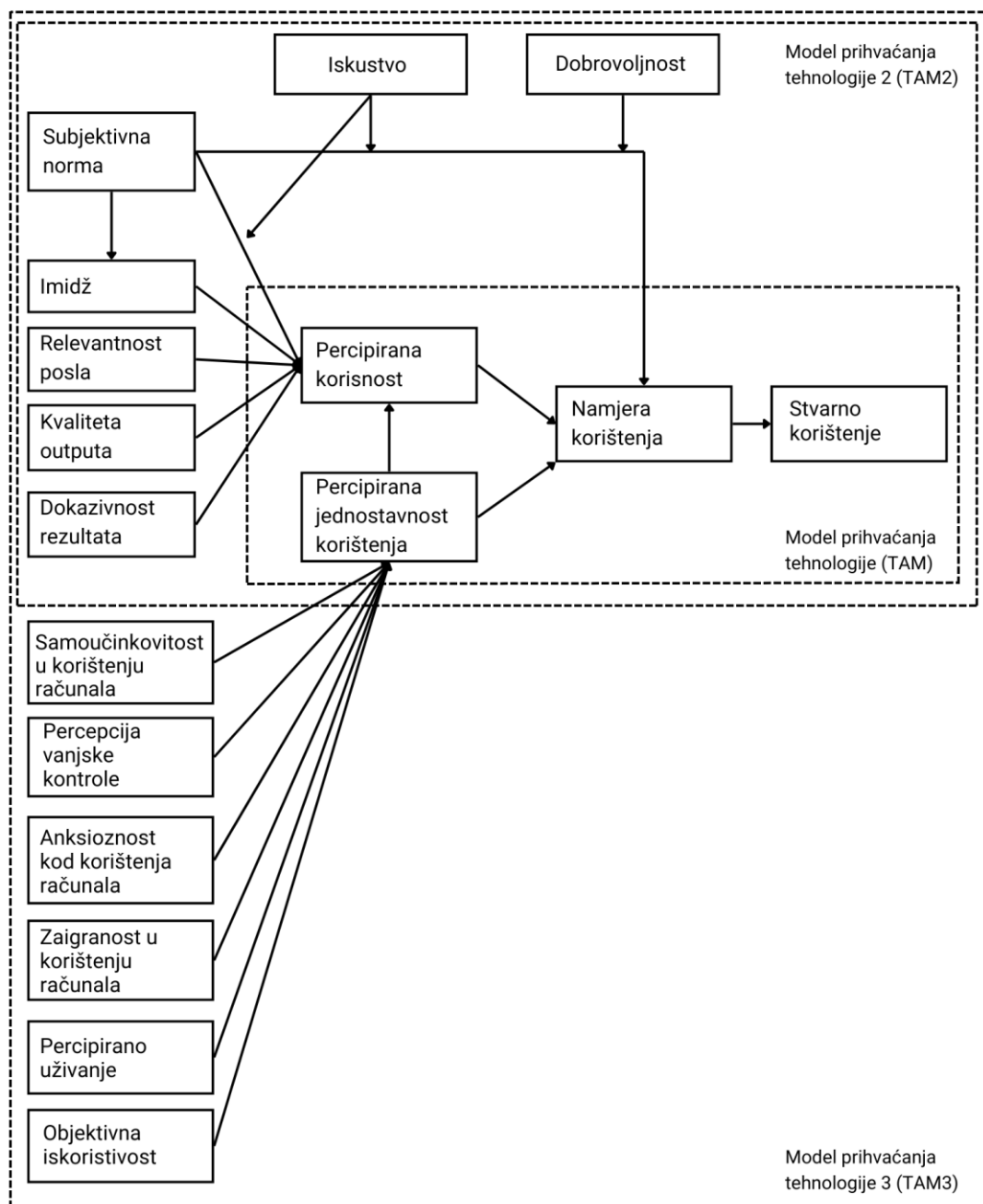


Slika 6. Prikaz modela – Model prihvaćanje tehnologije 2 (TAM2)

Izvor: prilagođeno prema: Venkatesh & Davis, 2000

Venkatesh i Bala (2008) proširuju TAM2, odnosno predstavljaju TAM3 koji je nastao tako što su se spojili TAM2 (Venkatesh & Davis, 2000) i Model determinanti percipirane jednostavnosti korištenja (Venkatesh, 2000). Model prihvaćanja tehnologije 3 dodatno proširuje model uključivanjem društvenog utjecaja, kognitivnih instrumenata procesa i emocionalnih procesa kako bi se bolje razumjelo prihvaćanje tehnologije i ponašanje pri korištenju (Venkatesh & Bala, 2008). S obzirom na to da je TAM3 predstavljen pet godina nakon UTAUT-a, on nije utjecao na formiranje izvornog UTAUT-a, ali jest na drugu verziju Ujedinjene teorije i prihvaćanja i korištenja tehnologije. TAM2 i TAM3 predstavljaju napredak u odnosu na izvorni TAM jer uzimaju u obzir dodatne čimbenike koji utječu na prihvaćanje i korištenje tehnologije

što model čini sveobuhvatnim i održava složenost prirode prihvatanja i korištenja tehnologije u različitim kontekstima.



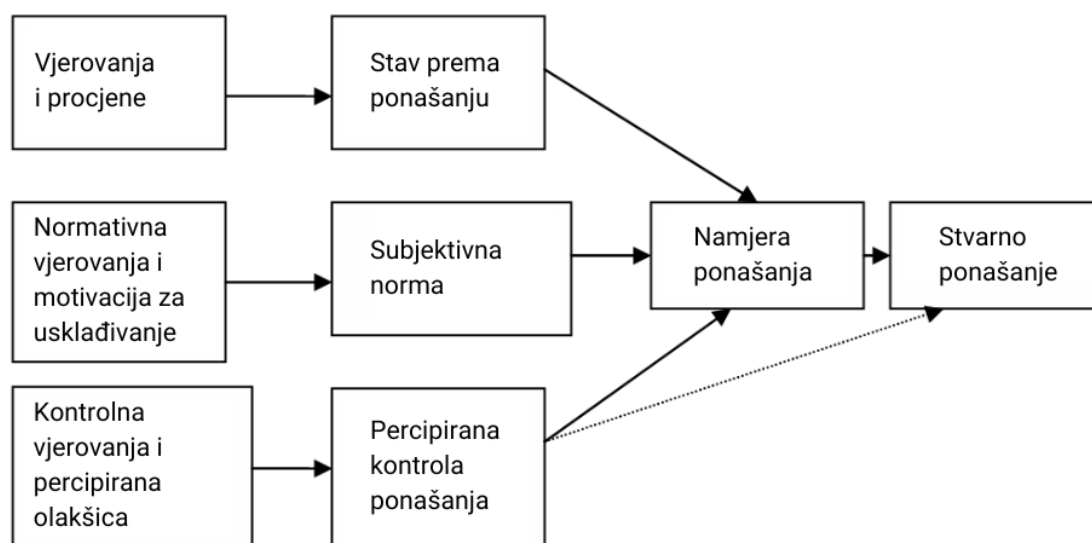
Slika 7. Prikaz razvoja modela TAM 3

Izvor: prilagođeno prema: Venkatesh & Bala, 2008

Motivacijski model (MM) vrlo je prihvaćen model u polju psihologije kao teorija koja dobro objašnjava nečije ponašanje, odnosno motivaciju da se tako ponaša. Motivacijski model može se koristiti u raznim kontekstima kako bi se provjerila motivacija ljudi da se ponašaju na određeni način te su tako Davis i sur. (1992) testirali motivacijski model u kontekstu upotrebe

tehnologije, odnosno računala, u poslovnom okruženju. Temeljni su konstrukti navedenog modela vanjska motivacija i unutarnja motivacija (Davis i sur., 1992; Venkatesh i sur., 2003). Davis i sur. (1992) definiraju vanjsku motivaciju kao želju korisnika da obavlja određenu aktivnost jer se ona percipira kao vrijednosni ishod različit od same aktivnosti, kao primjerice poboljšana radna učinkovitost, promaknuće ili financijska nagrada. S druge strane, unutarnja motivacija definira se kao percepcija da korisnici žele obavljati aktivnost bez očitih vrijednosnih ishoda osim samog procesa izvođenja aktivnosti (Venkatesh i sur., 2003).

Teorija planiranog ponašanja (TPB) je baš poput teorije razumnog djelovanja, utjecajni model iz domene razumijevanja ljudskog ponašanja, a vrlo često se primjenjuje u kontekstu prihvatanja i korištenja određene tehnologije. Teorija planiranog ponašanja zapravo je nadograđena Teorija razumnog djelovanja u koju je dodan konstrukt percipirane kontrole ponašanja koji se odnosi na percipiranu lakoću ili poteškoću izvođenja ponašanja, odnosno percepciju unutarnjih ili vanjskih ograničenja ponašanja u kontekstu istraživanja iz područja informacijskih sustava (Ajzen, 1991; Taylor & Todd, 1995a; Venkatesh i sur., 2003).

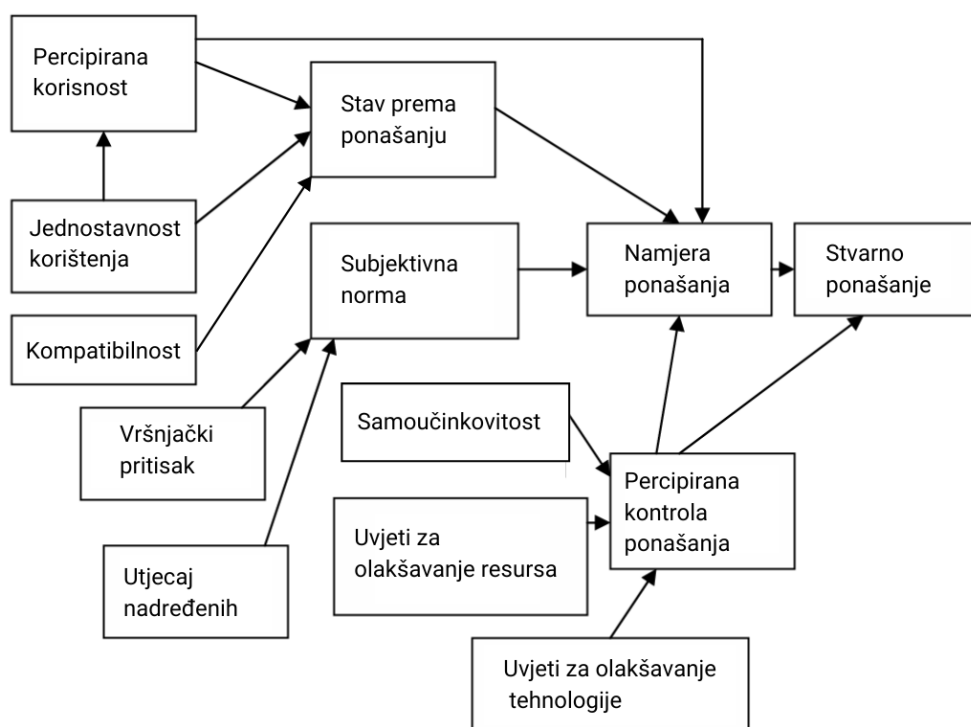


Slika 8. Prikaz modela – Teorija planiranog ponašanja TPB

Izvor: Ajzen, 1991

Kombinacija Modela prihvatanja tehnologije i Teorije planiranog ponašanja (C-TAM-TPB) hibridni je model koji ujedinjuje prediktore iz Teorije planiranog ponašanja sa percipiranom korisnošću iz Modela prihvatanja tehnologije, a koji su razvili Taylor i Todd (1995b). Temeljni konstrukti navedenog modela su stav prema ponašanju, subjektivna norma, percipirana kontrola ponašanja i percipirana korisnost. Ovaj model poznat je i kao rastavljena teorija planiranog

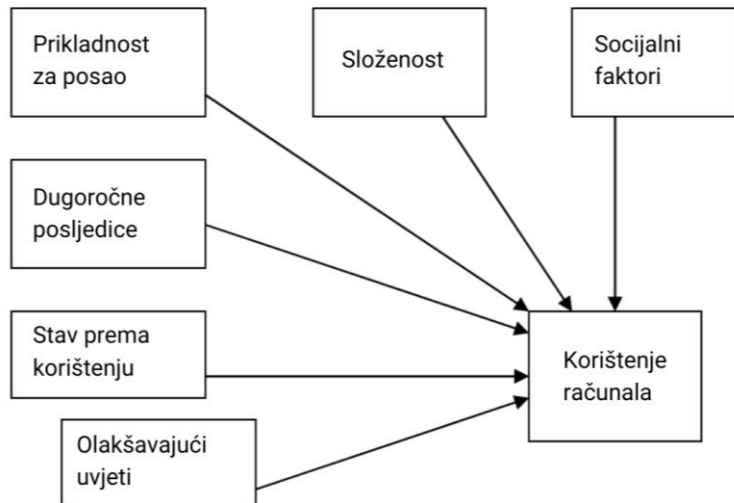
ponašanja (engl. *Decomposed theory of planned behavior*) jer je struktura vjerovanja u modelu rastavljena na komponente (Li, 2010).



Slika 9. Prikaz modela – kombinacija Modela prihvaćanja tehnologije i Teorije planiranog ponašanja (C-TAM-TPB)

Izvor: prilagođeno prema: Taylor & Todd, 1995b

Model korištenja osobnih računala izveden je najvećim dijelom iz Teorije ljudskog ponašanja (Triandis, 1977). Thompson i sur. (1991) prilagodili su Teoriju ljudskog ponašanja za istraživanja u kontekstu informacijskih sustava i koristili su ga za predviđanje korištenja osobnih računala s ciljem predviđanja ponašanja korištenja umjesto namjere. Temeljni konstrukti navedenog modela su bili: prikladnost za posao, složenost, dugoročne posljedice, utjecaj na korištenje, društveni čimbenici i olakšavajući uvjeti. Thompson i sur. (1991) navode kako je ponašanje određeno onim što ljudi žele raditi (stavovi), onim što ljudi misle da bi trebali raditi (socijalne norme), onim što obično rade (navike) te očekivanim posljedicama njihova ponašanja.



Slika 10. Prikaz modela – Model korištenja osobnih računala

Izvor: prilagođeno prema: Thompson i sur., 1991

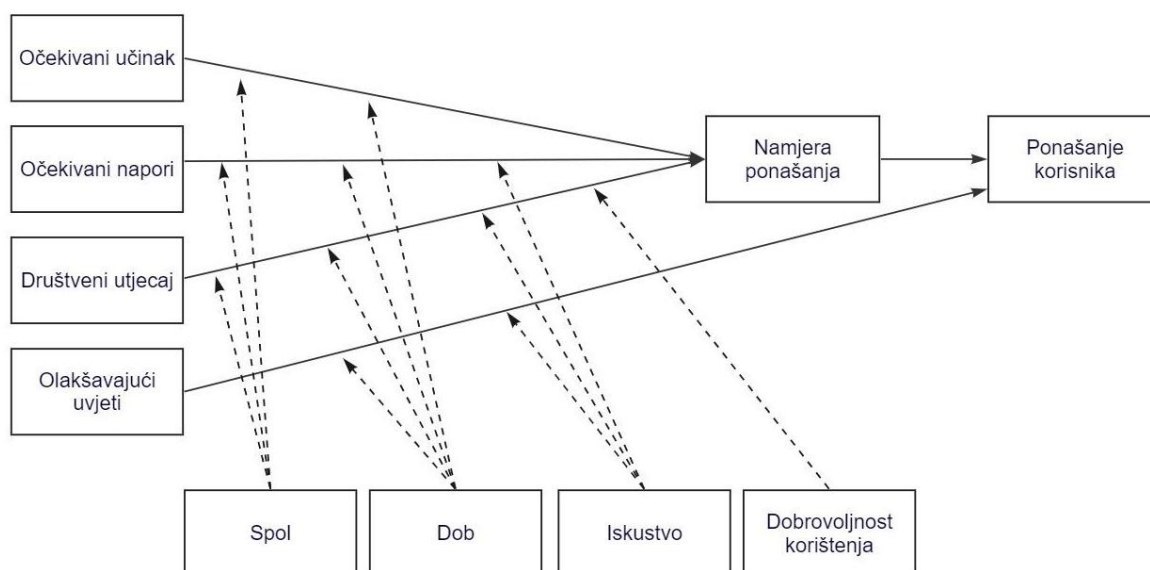
Teoriju difuzije inovacija (IDT) utemeljio je Rogers (1962) te se od tada u sociologiji koristila za proučavanje raznih inovacija. Unutar informacijskih sustava Moore i Benbasat (1991) prilagodili su karakteristike inovacija predstavljenih u Rogersovom radu i usavršili skup konstrukta koji bi se mogli koristiti za proučavanje individualnog prihvatanja tehnologije. Temeljni su konstrukti navedene teorije: relativna prednost, jednostavnost korištenja, slika, vidljivost, kompatibilnost, dokazivost rezultata i dobrovoljnost korištenja (Venkatesh i sur., 2003).

Socijalno-kognitivnu teoriju (SCT) ljudskog ponašanja utemeljio je Bandura (1986), a kasnije su Compeau i Higgins (1995) primijenili temelj navedene teorije i proširili na kontekst korištenja računala. U spomenutom proširenom modelu naglasak je bio na korištenju računala, ali je priroda modela takva da se može primijeniti i u kontekstu drugih tehnologija, a specifičan je po tome što korištenje koristi kao zavisnu varijablu. Temeljni su konstrukti navedenog modela: očekivani ishod izvedbe, osobni očekivani ishod, samoučinkovitost, utjecaj i anksioznost.

3.1.2. Konstrukti modela UTAUT

Na temelju navedenih modela nastajala je ujedinjena teorija prihvatanja i korištenja tehnologije (UTAUT). Navedeni model sadrži četiri prediktora: očekivani učinak (PE - *Performance expectancy*), očekivani naponi (EE - *Effort Expectancy*), društveni utjecaj (SI - *Social influence*) koji utječu na namjeru ponašanja te neizravno i na ponašanje korisnika te olakšavajući uvjeti

(FC - *Facilitating conditions*) koji su prediktor korištenja tehnologije. Također, model sadrži i četiri moderatora, a to su spol, dob, iskustvo i dobrovoljnost korištenja.



Slika 11. Prikaz UTAUT modela

Izvor: prilagođeno prema Venkateshu i sur. (2003)

3.1.2.1. PE – Očekivani učinak

Kako navode Venkatesh i sur. (2003) u izvornom radu, konstrukt očekivanog učinka ili očekivane izvedbe nastao je kao kombinacija pet prethodno korištenih konstrukta, a to su percipirana korisnost koja se koristila u modelima TAM i TAM2 te C-TAM-TPB modelu, vanjska motivacija koja se koristila u motivacijskom modelu (MM), usklađenost s poslom koji se koristio u MPCU modelu, konstrukt relativne prednosti koji se koristio u IDT modelu te konstrukt očekivanog ishoda korišten u SCT modelu. Venkatesh i sur. (2003) također navode da su čak i autori koji su razvijali navedene konstrukte priznavali sličnost među pojedinim konstruktima kao što su percipirana korisnost i usklađenost s poslom (Thompson i sur., 1991), percipirana korisnost i relativna prednost (Davis i sur., 1989; Moore i Benbasat, 1991; Plouffe i sur., 2001), percipirana korisnost i očekivani ishodi (Compeau i Higgins 1995), te usklađenost s poslom i očekivani ishodi (Compeau i Higgins, 1995).

Tablica 2. Korijeni konstrukta – Očekivani učinak (PE)

[PE] – Očekivani učinak – korijeni konstrukta, definicije i čestice konstrukta		
Konstrukt	Definicija	Čestice

Percipirana korisnost (Davis i sur., 1989)	Stupanj do kojeg osoba vjeruje da bi korištenje određenog sustava poboljšalo njegov ili njezin radni učinak	1. Korištenje sustava u mom poslu omogućilo bi mi da zadatke obavljam brže. 2. Korištenje sustava poboljšalo bi moju radnu učinkovitost. 3. Korištenje sustava u mom poslu povećalo bi moju produktivnost. 4. Korištenje sustava povećao bi moju učinkovitost na poslu. 5. Korištenje sustava olakšalo bi mi obavljanje posla. 6. Smatrao/la bih sustav korisnim u svom poslu.
Vanjska motivacija (Davis i sur., 1992)	Percepcija da će korisnici htjeti izvršiti neku aktivnost jer se smatra da je ključna u postizanju vrijednosnih ishoda koji se razlikuju od same aktivnosti, kao što je poboljšana radna učinkovitost, plaća ili napredovanja	Koristi iste čestice kao percipirana korisnost (TAM), navedeno iznad
Usklađenost s poslom (Thompson i sur., 1991)	Kako sposobnost sustava poboljšavaju individualnu radnu učinkovitost	1. Korištenje sustava neće imati učinka na moju radnu učinkovitost. 2. Korištenje sustava može smanjiti vrijeme potrebno za moje važne poslovne odgovornosti. 3. Korištenje sustava može značajno povećati kvalitetu mog rada. 4. Korištenje sustava može povećati učinkovitost izvršavanja poslovnih zadataka. 5. Korištenje može povećati količinu proizvodnje uz isti uloženi napor. 6. Uzimajući u obzir sve zadatke, opći dojam do mjere korištenje sustava može pomoći na poslu. *drugačija ljestvica za svaku česticu
Relativna prednost (Moore & Benbasat, 1991)	Mjera do koje se upotreba inovacije smatra boljom od upotrebe njezinog prethodnika	1. Korištenje sustava omogućava mi brže obavljanje zadataka. 2. Korištenje sustava poboljšava kvalitetu mog rada. 3. Korištenje sustava olakšava obavljanje mog posla.

		4. Korištenje sustava povećava moju učinkovitost na poslu. 5. Korištenje sustava povećava moju produktivnost.
Očekivani ishodi (Compeau i Higgins, 1995; Compeau i sur., 1999)	Očekivani ishodi odnose se na posljedice povezane s ponašanjem. Compeau i Higgins odvojili su poslovna očekivanja od individualnih osobnih očekivanja te su Venkatesh i sur. (2003) iskoristili 4 čestice očekivanja za posao i 3 čestice za osobna očekivanja.	1. Povećat ću svoju učinkovitost na poslu. 2. Provest ću manje vremena na rutinskim poslovnim zadacima. 3. Povećat ću kvalitetu ishoda mog posla. 4. Povećat ću količinu ishoda za istu količinu uloženog truda. 5. Moji suradnici će me smatrati kompetentnim. 6. Povećat ću svoje šanse za dobivanje promaknuća. 7. Povećat ću svoje šanse za dobivanje povišice.

Izvor: prilagođeno prema Venkateshu i sur., 2003

Shodno navedenim definicijama, konstrukt očekivanog učinka može se definirati kao stupanj u kojem pojedinac smatra da pojedina tehnologija unaprjeđuje njegovu učinkovitost i produktivnost ili pomaže u postizanju željenih ciljeva (Venkatesh i sur., 2003; Venkatesh i sur., 2012). Na temelju empirijskih podataka Venkatesh i sur. (2003) zaključuju kako je očekivani ishod najsnažniji prediktor namjere te ostaje značajan u svim točkama mjerenja, neovisno o tome je li riječ o mjerenjima u dobrovoljnim ili obveznim postavkama. Također, postoji i teorijska pretpostavka kako će na odnos između očekivanog učinka i namjere biti moderiran spolom i dobi zbog činjenice da su muškarci skloniji biti visoko usmjereni na zadatak (Minton i Schneider, 1980) te bi zbog toga očekivanja učinka, orijentiranih na zadatak, vjerojatno bila značajnija za muškarce. Slično spolu, teorija ukazuje i na pretpostavku da bi dob trebala imati moderirajući ulogu, a na to ukazuju i istraživanja vezana za stavove prema poslu koja sugeriraju da mlađi radnici pridaju više važnosti vanjskim nagradama (Hall & Mansfield, 1975; Porter, 1963). Također, Morris i Venkatesh (2000) prikazuju kako doista postoje razlike među spolovima i različitim dobnim skupinama kada je u pitanju usvajanje tehnologije, no Levy (1988) dodaje kako je pogrešno promatrati spol bez osvrta na dob.

Uzevši u obzir svu navedenu teoriju i provedeno empirijsko istraživanje sa svim ponuđenim česticama Venkatesh i sur. (2003) kreiraju konstrukt očekivanih ishoda čije su čestice sljedeće:

1. Smatrao bih sustav korisnim u svom poslu.

2. Korištenje sustava omogućava mi da zadatke obavljam brže.
3. Korištenje sustava povećava moju produktivnost.
4. Ako koristim sustav, povećat će se moje šanse za dobivanje povišice.

Ipak, Venkatesh i sur. (2012) u UTAUT2 modelu na istraživanju o prihvatanju i korištenju mobilnog interneta navode iduće čestice:

1. Smatram mobilni internet korisnim u svakodnevnom životu.
2. Korištenje mobilnog interneta povećava moje šanse za ostvarivanjem stvari koje su mi bitne.
3. Korištenje mobilnog interneta pomaže mi da brže obavim aktivnosti.
4. Korištenje mobilnog interneta povećava moju produktivnost.

Važno je za uočiti kako u kontekstu UTAUT2 modela Venkatesh i sur. (2012) ne koriste buduće vrijeme, već sadašnje te umjesto „Smatrao bih“ (engl. *I would find*) navode „Smatram“ (engl. *I find*)

U kontekstu ovih pitanja korišten je termin sustava premda se umjesto njega može ubacivati naziv tehnologije u kontekstu kojeg se provodi istraživanje. Nadalje, Venkatesh i sur. (2003) kreiraju i hipotezu vezanu za konstrukt očekivanih ishoda koja glasi: *Utjecaj očekivanog učinka na namjeru ponašanja bit će moderiran spolom i dobi, tako da će efekt biti jači za muškarce, a posebno za mlađe muškarce.*

Nadalje, zanimljivo je spomenuti kako Venkatesh i sur. (2012) navode da je čestica „Korištenje mobilnog interneta povećava moje šanse za ostvarivanjem stvari koje su mi bitne.“ izostavljena u konačnoj analizi jer se nije pokazala važnom. Ipak, druge studije koje istražuju prihvatanje tehnologije u pravilu ipak ostavljaju ovaj konstrukt sa sve četiri navedene čestice s ciljem bolje usporedivosti s ovim i srodnim istraživanjima.

3.1.2.2. EE – Očekivani naponi

Venkatesh i sur. (2003) navode kako se konstrukt očekivanih napora definira kao stupanj lakoće povezan s korištenjem nekakvog sustava ili tehnologije. Navedeni konstrukt nastao je kao kombinacija percipirane lakoće korištenja iz TAM i TAM2 modela, kompleksnosti iz MPCU modela te lakoće korištenja iz IDT modela. Kao i za konstrukt očekivanog učinka, za konstrukt očekivanih napora u teoriji se i ranije pisalo o sličnosti među konstruktima koji su prethodno navedeni (Davis i sur., 1989; Moore i Benbasat, 1991; Thompson i sur., 1991).

Tablica 3. Korijen konstrukta – Očekivani naponi (EE)

[EE] – Očekivani naponi – korijeni konstrukta, definicije i čestice konstrukta		
Konstrukt	Definicija	Čestice
Percepcija jednostavnosti korištenja (Davis i sur., 1989)	Stupanj do kojeg osoba vjeruje da će korištenje sustava biti bez napora	1. Učenje za upravljanje sustavom bilo bi mi lako. 2. Bilo bi mi lako dobiti sustav da radi ono što ja želim. 3. Moja interakcija sa sustavom bila bi jasna i razumljiva. 4. Sustav bi mi bio fleksibilan za interakciju. 5. Bilo bi mi lako postati vješt u korištenju sustava. 6. Smatrao/la bih sustav lakim za korištenje.
Kompleksnost (Thompson i sur., 1991)	Stupanj do kojeg se sustav percipira kao relativno težak za razumijevanje	1. Korištenje sustava uzelo bi mi previše vremena od mojih normalnih dužnosti. 2. Rad s ovim sustavom je toliko kompliciran da je teško razumjeti što se događa. 3. Korištenje sustava traži previše vremena provedenog u mehaničkim operacijama (npr. unos podataka). 4. Trebalo bi previše vremena da naučim koristiti sustav da bi bilo vrijedno truda.
Jednostavnost korištenja (Moore i Benbasat, 1991)	Stupanj do kojeg se inovacija u korištenju percipira kao teška za korištenje	1. Moja interakcija sa sustavom je jasna i razumljiva. 2. Vjerujem da je lako dobiti sustav da ono što ja želim. 3. Općenito, vjerujem da je sustav lako koristiti.

		4. Učenje za upravljanje sustavom mi je lako.
--	--	---

Izvor: prilagođeno prema Venkateshu i sur. (2003)

Uzevši u obzir svu navedenu teoriju i provedeno empirijsko istraživanje sa svim ponuđenim česticama Venkatesh i sur. (2003) kreiraju konstrukt očekivanih napora čije su čestice sljedeće:

1. Moja interakcija sa sustavom bila bi jasna i razumljiva.
2. Bilo bi mi lako postati vješt/a u korištenju sustava.
3. Smatrao/la bih sustav jednostavnim za korištenje.
4. Učenje za upravljanje sustavom je lako za mene.

Očekuje se da će konstrukti usmjereni na napor biti značajniji u ranoj fazi novog ponašanja, a s vremenom će važnost opadati do razine beznačajnosti (Venkatesh i sur., 2003; Agarwal i Prasad, 1997; Davis i sur., 1989). Venkatesh i Morris (2000), oslanjajući se na istraživanja Bem i Allen (1974) te Bozionelos (1996), navode kako će očekivani napor kod korištenja nekog sustava ili tehnologije biti izraženiji kod žena nego kod muškaraca, ali da jednako tako i dob osobe igra ulogu gdje istraživanja ukazuju na to da starija dob ima više izazova u obradi složenih podražaja i usmjeravanja pažnje na informacije na poslu (Plude i Hoyer, 1985; Venkatesh i sur., 2003).

Ipak, kao i kao čestica očekivanog učinka, kod pitanja UTAUT2 modela vidljivo je kako su tvrdnje unutar konstrukta postavljene u sadašnjem, a ne u budućem vremenu kako je to slučaj kod čestica unutar UTAUT-a pa tako one u glase (Venkatesh i sur., 2012):

1. Lako mi je naučiti koristiti mobilni internet.
2. Moja interakcija s mobilnim internetom je jasna i razumljiva.
3. Smatram da je mobilni internet jednostavan za korištenje.
4. Lako mi je postati vješt u korištenju mobilnog interneta.

U konačnici, Venkatesh i sur. (2003) kreiraju i hipotezu koja glasi: *Utjecaj očekivanih napora na namjeru ponašanja bit će moderiran spolom, dobi i iskustvom, tako da će efekt biti jači za žene, posebno mlađe žene i posebno u ranoj fazi iskustva.*

3.1.2.3. SI – Društveni utjecaj

Konstrukt društvenog utjecaja definira se kao stupanj do kojeg pojedinac percipira da važne osobe iz okruženja smatraju da on/ona trebaju koristiti novi sustavi ili tehnologiju (Venkatesh i sur., 2003). Konstrukt društvenog utjecaja vrlo je srodan konstruktu subjektivne norme

predstavljenom u modelima TRA, TAM2, TPB i C-TAM-TPB, ali u korijenu konstrukta društvenog utjecaja su i društveni faktori iz modela MPCU, kao i imidž iz IDT modela (Venkatesh i sur., 2003). Empirijski podatci pokazuju kako se svi navedeni konstrukti ponašaju relativno slično jer niti jedan od njih nije značajan u dobrovoljnim kontekstima, ali svi su značajni kada je upotreba pojedinog sustava ili tehnologije obavezna (Venkatesh i sur., 2003).

Tablica 4. Korijen konstrukta – društveni utjecaj (SI)

[SI] – Društveni utjecaj – korijeni konstrukta, definicije i čestice konstrukta		
Konstrukt	Definicija	Čestice
Subjektivna norma (Ajzen, 1991; Davis i sur., 1989; Fishbein & Ajzen, 1975; Taylor i Todd, 1995a)	Percepcija osobe da većina ljudi koji su joj važni misle da ona treba napraviti ono što je kontekst pitanja	1. Ljudi koji utječu na moje ponašanje misle da bih trebao/la koristiti sustav. 2. Ljudi koji su mi važni misle da bih trebao/la koristiti sustav.
Društveni faktori (Thompson i sur., 1991)	Individualizacija referentne grupe subjektivne kulture i specifični interpersonalni dogovori koje pojedinac ima s drugima u specifičnim društvenim situacijama	1. Koristim sustav jer to čine osobe iz moje referentne grupe. 2. Viši menadžment ove tvrtke bio je od pomoći u korištenju sustava. 3. Moj nadređeni je vrlo podržavajući u korištenju sustava za moj posao. 4. Općenito, organizacija podržava korištenje sustava.
Imidž (Moore & Benbasat, 1991)	Stupanj do kojeg korištenje inovacije povećava percepciju osobe o vlastitom imidžu ili statusu u društvenom sustavu	1. Ljudi u mojoj organizaciji koji koriste sustav imaju veći status. 2. Ljudi u mojoj organizaciji koji sustav su visoko profilirani. 3. Imati sustav je statusni simbol u mojoj organizaciji.

Izvor: prilagođeno prema Venkateshu i sur. (2003)

Uzevši u obzir svu navedenu teoriju i provedeno empirijsko istraživanje sa svim ponuđenim česticama Venkatesh i sur. (2003) kreiraju konstrukt društvenog utjecaja čije su čestice sljedeće:

1. Ljudi koji utječu na moje ponašanje misle da bih trebao/la koristiti sustav.
2. Ljudi koji su mi važni misle da bih trebao/la koristiti sustav.
3. Viši menadžment ove tvrtke bio je od pomoći u korištenju sustava.
4. Općenito, organizacija podržava korištenje sustava.

Uloga društvenog utjecaja o prihvatanju tehnologije vrlo je kompleksna i podložna širokom rasponu uvjetovanih utjecaja, a utjecaj na ponašanje pojedinca može se promatrati kroz tri mehanizma: usklađenost, internalizaciju i identifikaciju (Venkatesh i Davis, 2000; Venkatesh i sur., 2003). Mehanizmi internalizacije i identifikacije odnose se na strukturu stavova pojedinca s ciljem da reagira na potencijalne prednosti korištenja tehnologije kao što je bolji društveni status, dok mehanizam usklađenosti za cilj ima promjenu namjere upravo zbog društvenog pritiska, odnosno cilj ovog mehanizma je da se pojedinac uskladi sa društvenim očekivanjima (Venkatesh i sur., 2003). Teorija upućuje na to da su žene sklonije biti osjetljivije na mišljenja drugih ljudi i stoga smatraju da je društveni utjecaj značajan pri formiranju namjere za korištenje nove tehnologije, premda učinak opada s iskustvom (Miller, 1976; Venkatesh i sur., 2000; Venkatesh i Morris, 2000). Također, Rhodes (1983) navodi kako se potrebe za pripadanjem i društvenim prihvatanjem povećavaju s godinama, sugerirajući da osobe starije dobi pridaju veću važnost društvenim utjecajima, ali i da taj utjecaj opada s iskustvom (Venkatesh i Morris, 2000). U konačnici, s obzirom na navedenu teoriju, Venkatesh i sur. (2003) formiraju hipotezu koja glasi: *Utjecaj društvenog utjecaja na namjeru ponašanja bit će moderiran spolom, dobi, dobrovoljnošću i iskustvom, tako da će efekt biti jači za žene, posebno starije žene, posebno u obveznim postavkama u ranoj fazi iskustva.*

Ipak, u odnosu na UTAUT, Venkatesh i sur. (2012) u istraživanju o prihvatanju i korištenju mobilnog interneta, u kojem su predstavili UTAUT2 model, navode nešto drugačije čestice unutar konstrukta društvenog utjecaja:

1. Osobe koje su mi važne misle da bih trebao koristiti mobilni internet.
2. Osobe koje utječu na moje ponašanje misle da bih trebao koristiti mobilni internet.
3. Osobe čije mišljenje cijenim preferiraju da koristim mobilni internet.

S obzirom na to da je konstrukt društvenog utjecaja unutar UTAUT-a orijentiran na društveni utjecaj unutar organizacije u kojoj pojedinac radi, čestice tog konstrukta ne poklapaju se s česticama unutar UTAUT2 modela jer je on orijentiran na prihvatanje i korištenje tehnologije kod krajnjih korisnika u neorganizacijskim postavkama. Iz istog je razloga moderator dobrovoljnosti izbačen iz modela pa tako i iz ranije navedene hipoteze što u konačnici dovodi

do dorađene verzije hipoteze koja glasi: *Utjecaj društvenog utjecaja na namjeru ponašanja bit će moderiran spolom, dobi i iskustvom, tako da će efekt biti jači za žene, posebno starije žene, posebno u ranoj fazi iskustva.*

3.1.2.4. FC – Olakšavajući uvjeti

Konstrukt olakšavajućih uvjeta definira se kao stupanj do kojeg pojedinac vjeruje da postoji organizacijska i tehnička infrastruktura koja podržava korištenje sustava, a ova definicija utjelovljuje koncepte iz tri različita konstrukta: percipirana bihevioralna kontrola prisutna u modelima TPB i C-TAM-TPB, uvjeti olakšavanja prisutan u modelu MPCU te kompatibilnost prisutna u IDT modelu (Venkatesh i sur., 2003). Navedeni konstrukti operacionalizirani su tako da uključuju aspekte tehnološkog i/ili organizacijskog okruženja koji su dizajnirani kako bi uklonili prepreke za upotrebu, a empirijski dokazi navedeni u radu Venkatesha i sur. (2003) sugeriraju da su odnosi među navedenim konstruktima vrlo slični.

Tablica 5. Korijeni konstrukta – olakšavajući uvjeti (FC)

[FC] – Olakšavajući uvjeti – korijeni konstrukta, definicije i čestice konstrukta		
Konstrukt	Definicija	Čestice
Percepcija ponašanja (Ajzen 1991; Taylor i Todd 1995a)	Održava unutarnje i vanjske ograničavajuće faktore te kontrole ponašanja uključujući samoučinkovitost, olakšavanje resursa i uvjete koji omogućuju tehnologiju	1. Imam kontrolu nad korištenjem sustava. 2. Imam resurse potrebne da koristim sustav. 3. Imam znanje neophodno da koristim sustav. 4. S obzirom na resurse, prilike i znanje koje imam, bilo bi mi lako koristiti sustav. 5. Sustav nije kompatibilan s drugim sustavima koje koristim.
Uvjeti olakšavanja (Thompson i sur., 1991)	Objektivni faktori u okruženju koji čine da se promatrači slažu kako je neki čin lako učiniti, uključujući mogućnost pružanja podrške	1. Upute o upravljanju su mi dostupne pri odabiru sustava. 2. Specijalizirane upute o upravljanju sustavom su mi dostupne.

		3. Određena osoba ili grupa ljudi su mi na raspolaganju za pomoć s poteškoćama u upravljanju sustavom.
Kompatibilnost (Moore & Benbasat, 1991)	Stupanj do kojeg se inovacija percipira kao usklađena s postojećim uvjerenjima, prošlim iskustvima i potrebama potencijalnih korisnika	1. Mislim da je sustav kompatibilan sa svim aspektima mog rada. 2. Mislim da bi korištenje sustava bilo dobro u skladu s načinom na koji radim. 3. Korištenje sustava uklapa se u moj stil rada.

Izvor: prilagođeno prema Venkateshu i sur. (2003)

Uzevši u obzir svu navedenu teoriju i provedeno empirijsko istraživanje sa svim ponuđenim česticama Venkatesh i sur. (2003) kreiraju konstrukt olakšavajućih uvjeta čije su čestice sljedeće:

1. Imam potrebne resurse za korištenje sustava.
2. Imam potrebno znanje za korištenje sustava.
3. Sustav nije kompatibilan s drugim sustavima koje koristim.
4. Određena osoba (ili grupa osoba) dostupna je za pomoć pri poteškoćama u upravljanju sustavom.

Venkatesh i sur. (2012) u UTAUT2 modelu navode jednu ključnu razliku kod čestica olakšavajućih uvjeta, a to je da čestica pod rednim brojem tri nije u negaciji, pa tako čestice unutar UTAUT2 modela glase:

1. Imam potrebne resurse za korištenje mobilnog interneta.
2. Imam potrebno znanje za korištenje mobilnog interneta.
3. Mobilni internet kompatibilan je s drugim tehnologijama koje koristim.
4. Mogu dobiti pomoć od drugih kada imam poteškoća s korištenjem mobilnog interneta.

Empirijski dokazi koje navode Venkatesh i sur. (2003) sugeriraju kako su problemi vezani za infrastrukturnu podršku ključni koncept unutar konstrukta očekivanja napora koji ispituje lakoću primjene određenog sustava ili tehnologije. Isti autori također navode kako očekivani napori nisu značajni ukoliko su u modelu prisutni očekivani učinak i konstrukt očekivanih napora zbog

čega autori predlažu prvi dio hipoteze: *Olakšavajući uvjeti neće imati značajan utjecaj na namjeru ponašanja.*

Ipak, Venkatesh i sur. (2003) navode kako njihovi empirijski podaci upućuju na to da olakšavajući uvjeti imaju izravan utjecaj na korištenje, odnosno konstrukt ponašanja korisnika. Navedeni autori ukazuju na to da će spomenuti utjecaj rasti s iskustvom jer korisnici imaju razne potrebe da zatraže pomoć s upravljanjem određene tehnologije. Psiholozi upućuju na to da je mogućnost da zatraže pomoć nekoga od kolega s posla posebno bitna starijim radnicima (Hall & Mansfield, 1975). Na temelju navedene teorije i dobivenih rezultata istraživanja Venkatesha i Morrisa (2000), Venkatesh i sur. (2003) u izvornom radu kod izrade Ujedinjene teorije prihvaćanja i korištenja tehnologije predlažu drugi dio hipoteze vezan za olakšavajuće uvjete, a koji glasi: *Utjecaj olakšavajućih uvjeta na korištenje bit će moderiran dobi i iskustvom, na način da će efekt biti jači za starije ispitanike, posebno s povećanjem iskustva.*

Izvorni rad u kojem je predstavljena Ujedinjena teorija prihvaćanja i korištenja tehnologije uključuje i konstrukt anksioznosti u vezi računala te konstrukt samoučinkovitosti jer unutar Socijalno-kognitivne teorije (SCT) oni predstavljaju značajne prediktore namjere, no bitno je napomenuti kako SCT nema konstrukt očekivanih napora. S druge strane, UTAUT ne uključuje anksioznost i samoučinkovitost kao izravne prediktore jer je Venkatesh (2000) ukazao na to da su samoučinkovitost i anksioznost konceptualno i empirijski različiti od očekivanih napora, odnosno percipirane lakoće korištenja (Venkatesh i sur., 2003).

Osim spomenute samoučinkovitosti i anksioznosti, Venkatesh i sur. (2003) u izvornom su radu ispitali i konstrukt stava prema korištenju, a koji vuče korijene iz: stava prema ponašanju (TRA, TPB/DTPB, C-TAM-TPB), intrinzične motivacije (MM) i afekta (SCT). Konstrukti stavova u određenim slučajevima snažan su prediktor namjere (TRA, TPB/DTPB, MM), dok u drugim slučajevima konstrukt uopće nije značajan (TAM-TPB, MPCU, SCT). Empirijski dokazi upućuju na to da su konstrukti stavova značajni ukoliko konstrukti očekivanog učinka i očekivanih napora nisu uključeni u model. Iz navedenog razloga Venkatesh i sur. (2003) u izvornom radu o Ujedinjenoj teoriji prihvaćanja i korištenja tehnologije ne predviđaju stav prema korištenju tehnologije izravnim i značajnim prediktorom namjere korištenja tehnologije.

Čestice navedenih konstrukta su sljedeće:

Anksioznost:

1. Osjećam nelagodu zbog korištenja sustava.

2. Plaši me pomisao da bih mogao izgubiti puno informacija pritiskom na pogrešku tipku.
3. Oklijevam koristiti sustav zbog straha od pravljenja grešaka koje ne mogu ispraviti.
4. Sustav me donekle plaši.

Samoučinkovitost:

- Mogao bih učiniti radnju u sustavu:

1. i kada ne bih imao nikoga oko sebe tko bi mi govorio što trebam raditi.
2. i kada bih mogao nazvati nekoga za pomoć ako zapnem.
3. kada bih imao puno vremena dovršiti posao za koji je softver namijenjen.
4. kada bih imao samo ugrađenu pomoćnu funkciju za pomoć.

Stav prema korištenju tehnologije:

1. Korištenje sustava je loša/dobra ideja.
2. Sustav čini rad zanimljivim.
3. Rad sa sustavom je zabavan.
4. Volim raditi sa sustavom.

3.2. UTAUT2

Venkatesh i sur. (2012) proširuju Ujedinjenu teoriju prihvatanja i korištenja tehnologije sa tri dodatna konstrukta te ga nazivaju UTAUT2 (akronim od *Unified theory of accepting and use of technology 2*). U navedenom radu navodi se kako se izvorni UTAUT proširuje konstruktima: hedonistička motivacija (HM – engl. *Hedonic motivation*), vrijednost cijene (PV – engl. *Price value*) i navika (HT – engl. *Habit*), ali se i moderator dobrovoljnosti izbacuje iz modela. Također, UTAUT2 model ipak predviđa kako je konstrukt olakšavajućih uvjeta značajan prediktor namjere ponašanja. Navedena proširena verzija rezultirala je značajnim poboljšanjem objašnjenja varijance kod namjere korištenja sa 54 % na 74 % te povećanje objašnjenja varijance kod korištenja tehnologije sa 40 % na 52 %.

3.2.1. Konstrukti modela UTAUT2

Izvorni UTAUT uvelike je bio korišten u mnogobrojnim istraživanjima kao osnovni model u proučavanjima raznih tehnologija u organizacijskim i neorganizacijskim postavkama. Venkatesh i sur. (2012) navode kako postoje tri široke vrste proširivanja UTAUT-a, a to su proširenja u novim kontekstima, npr. novim tehnologijama, novim populacijama korisnika i

novim kulturnim postavkama. Druga vrsta proširivanja je dodavanje novih konstrukta kako bi se proširio opseg endogenih teorijskih mehanizama navedenih u izvornom UTAUT-u, a treća vrsta uključuje dodavanje egzogenih prediktora UTAUT varijabli. Navedena proširenja i integracije izvornog UTAUT-a poslužili su Venkateshu i sur. (2012) za konstrukciju UTAUT2 modela. Ono što je temeljna razlika između UTAUT-a i UTAUT2 modela jest činjenica da je UTAUT izvorno razvijen za objašnjavanje prihvatanja i korištenja tehnologije u organizacijskom okruženju, odnosno u organizacijskim postavkama, dok se UTAUT2 usmjerava na prihvatanje i korištenje tehnologije kod krajnjih potrošača.

3.2.1.1. *HM – Hedonistička motivacija*

Venkatesh i sur. (2012) kao prvi od navedena tri novododana konstrukta navode hedonističku motivaciju iz razloga što mnogi autori poput Holbrooka i Hirschmana (1982), Nysveena i sur. (2005), van der Heijdena (2004) te Browna i Venkatesha (2005) navode upravo uživanje kao važan prediktor uporabe tehnologije od strane krajnjih potrošača u neorganizacijskim postavkama. Također, Venkatesh i sur. (2012) naglašavaju kako je glavni razlog integracije hedonističke motivacije u model zapravo činjenica da on nadopunjuje korisnost, odnosno očekivani učinak, najjači prediktor izvornog UTAUT-a. Nadalje, s obzirom na to da je UTAUT2 usmjeren na krajnje potrošače, a ne korisnike tehnologije unutar organizacije, njima je bitna i cijena proizvoda ili usluge jer su sami odgovorni za troškove, a troškovi ponekad dominiraju kod odluke o usvajanju neke tehnologije ili kupnje nekog proizvoda (Venkatesh i Brown, 2005; Chan i sur., 2008; Coulter i Coulter, 2007). Dakle, dodavanje konstrukta povezanog sa cijenom, odnosno monetarnim troškovima, nadopunit će postojeći konstrukt očekivanih napora koji se bazira na drugoj vrsti resursa. U konačnici, s obzirom na to da pojedini radovi osporavaju namjeru ponašanja kao prediktora korištenju tehnologije, Venkatesh i sur. (2012) uvode konstrukt navike kao prediktora namjere, ali i korištenja tehnologije (Limayen i sur., 2007; Davis i Venkatesh, 2004; Kim i sur., 2005). Venkatesh i sur. (2012) napominju kako je konstrukt navike integriran u UTAUT2 jer ima izravan utjecaj na korištenje tehnologije i oslabljuje ili ograničava snagu odnosa između namjere ponašanja i korištenja tehnologije. Navedeni autori napominju i kako je osim tri novododana konstrukta izbačena dobrovoljnost kao moderator, a olakšavajući uvjeti su osim korištenja tehnologije postali prediktor i namjeri ponašanja.

Venkatesh i sur. (2012) navode kako je konstrukt hedonističke motivacije preuzet i prilagođen prema Kim i sur. (2005). Hedonistička motivacija definirana je kao zabava ili užitek proizašao iz korištenja tehnologije, a pokazalo se kako igra važnu ulogu u određivanju prihvatanja i

uporabe tehnologije (Venkatesh i sur., 2012; Marikyan i Papagiannidis, 2023). Kako je vidljivo u tablici 6., Kim i sur. (2005) unutar konstrukta imaju četiri čestice te ga nazivaju hedonistička vrijednost (*engl. Hedonic value*), dok Venkatesh i sur. (2012) navode tri čestice, a sam konstrukt nazivaju hedonistička motivacija. Unutar tablice 6. čestice su navedene i na engleskom jeziku kako je izvorno navedeno u radu jer su navedeni termini koji predstavljaju hedonističke vrijednosti izazov za prevesti na hrvatski jezik zbog izrazite međusobne sličnosti. Korijen ovog konstrukta predstavili su Davis i sur. (1992) kada su imali konstrukt ugone/užitka (*engl. Enjoyment*), koji je također mjeren na sedmostupanjskoj Likertovoj ljestvici kao i ovaj konstrukt naveden u tablici 6., no Kim i sur. (2005) djelomično su modificirali navedeni konstrukt pa su Venkatesh i sur. (2012) po toj verziji prilagodili svoj istraživački instrument.

Tablica 6. Korijen konstrukta – hedonističke motivacije (HM)

Kim i sur. (2005)	Venkatesh i sur. (2012)
Korištenje ovog mrežnog sjedišta je zabavno. (<i>engl. Using this website is fun</i>)	Korištenje mobilnog interneta je zabavno. (<i>engl. Using mobile internet is fun</i>)
Korištenje ovog web-sjedišta mi je užitak. (<i>engl. Using this website is a joy to me</i>)	-
Korištenje ovog web-sjedišta je ugodno. (<i>engl. Using this website is enjoyable</i>)	Korištenje mobilnog interneta je ugodno. (<i>engl. Using mobile internet is enjoyable</i>)
Korištenje ovog web-sjedišta je vrlo zanimljivo. (<i>engl. Using this website is very entertaining</i>)	Korištenje mobilnog interneta vrlo je zanimljivo. (<i>engl. Using mobile internet is very entertaining</i>)

Izvor: prilagođeno prema – Kim i sur., 2005; Venkatesh i sur., 2012

3.2.1.2. PV – Vrijednost cijene

Kako navode Venkatesh i sur. (2012), konstrukt cijene (PV) ili vrijednosti cijene preuzet je i prilagođen prema Dodds i sur. (1991). Venkatesh i sur. (2012) svoj konstrukt nazivaju vrijednost cijene (*engl. Price value*), a Dodds i sur. (1991), od kojih je konstrukt preuzet i modificiran, konstrukt su nazvali percipirana vrijednost. Konstrukt vrijednost cijene definiran je kao kompromis potrošača između percipirane koristi tehnologije i novčanog troška za njezino korištenje (Venkatesh i sur., 2012; Marikyan i Papagiannidis, 2023). Venkatesh i sur. (2012) u smislu mjerenja od ispitanika traže da se na Likertovoj ljestvici od 1 do 7 izjasne koliko se slažu

s tvrdnjom, dok Dodds i sur. (1991) koriste semantički diferencijal te od ispitanika traže da se izraze od negativne vrijednosti do pozitivne vrijednosti, primjerice: dobra vrijednost za novac do loša vrijednost za novac ili pak vrlo prihvatljivo do vrlo neprihvatljivo.

Tablica 7. Korijen konstrukta – vrijednosti cijene (PV)

Dodds i sur. (1991)	Venkatesh i sur. (2012)
Proizvod je (<i>jako dobra vrijednost za novac – jako loša vrijednost za novac</i>).	Mobilni internet je dobre vrijednosti za novac.
Po ovoj cijeni proizvod je (<i>vrlo ekonomičan - vrlo neekonomičan</i>).	Po ovoj cijeni, mobilni internet pruža dobru vrijednost za novac.
Proizvod se smatra dobrom kupnjom.	-
Iskazana cijena za proizvod je (<i>vrlo prihvatljiva – vrlo neprihvatljiva</i>).	Mobilni internet je razumne cijene.
Čini se kako je cijena proizvoda prepovoljna.	-

Izvor: prilagođeno prema – Dodds i sur., 1991; Venkatesh i sur., 2012

3.2.1.3. HT – Navika

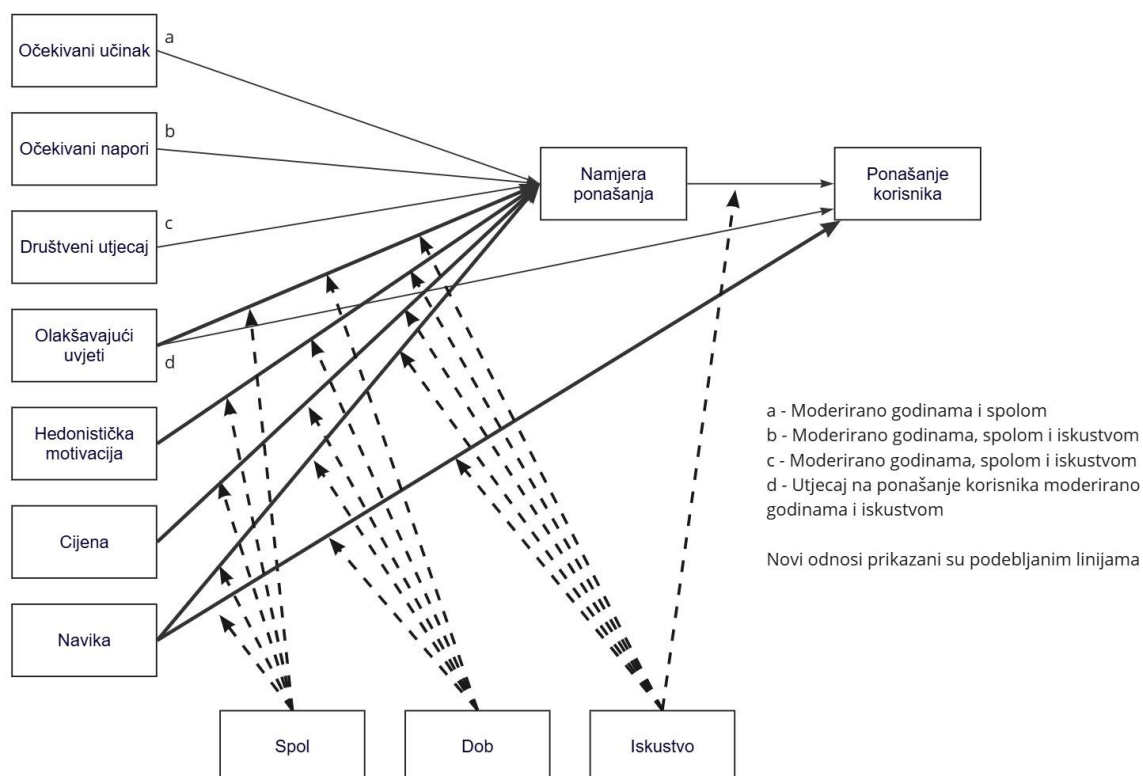
Venkatesh i sur. (2012) navode kako je konstrukt navike preuzet i prilagođen iz rada Limayen i Hirt (2003) koji su istraživali korištenje WebBoarda, internetske elektroničke oglasne ploče sa strukturom mapa. Naime, kako je vidljivo u tablici 7., Venkatesh i sur. (2012) izbacili su jednu česticu iz konstrukta. Nadalje, Limayen i Hirt (2003) koristili su Likertovu ljestvicu sa pet stupnjeva dok su Venkatesh i sur. (2012) koristili sa sedam stupnjeva. Konstrukt navike definira se kao stupanj do kojeg ljudi mogu imati tendenciju automatski izvoditi određena ponašanja (Venkatesh i sur., 2012; Marikyan i Papagiannidis, 2023).

Tablica 8. Korijen konstrukta – navika (HT)

Limayen i Hirt (2003)	Venkatesh i sur. (2012)
Korištenje WebBoarda postalo mi je navika.	Korištenje mobilnog interneta postalo mi je prirodno.
Ovisan sam o korištenju WebBoarda.	Ovisan sam o korištenju mobilnog interneta.
Moram koristiti WebBoard.	Moram koristiti mobilni internet.

Ne razmislim dvaput prije nego upotrijebim WebBoard.	-
Korištenje WebBoarda postalo mi je prirodno.	Korištenje mobilnog interneta postalo mi je prirodno.

Izvor: prilagođeno prema – Limayen i Hirt, 2003; Venkatesh i sur., 2012



Slika 12. Prikaz modela UTAUT2

Izvor: prilagođeno prema Venkateshu i sur. (2012)

3.3. Promjene unutar UTAUT2 modela u odnosu na izvorni UTAUT

Jedna od prvih uvedenih promjena UTAUT2 modela u odnosu na izvorni UTAUT je dodavanje izravnog odnosa olakšavajućih uvjeta na namjeru ponašanja jer se UTAUT2 usredotočuje na korištenje tehnologije od strane krajnjih potrošača. Ajzen (1991) navodi kako se pretpostavlja da će olakšavajući uvjeti (FC) izravno utjecati na korištenje tehnologije (USE) na temelju ideje da će u organizacijskom okruženju olakšavajući uvjeti služiti kao zamjena za stvarnu bihevioralnu kontrolu i izravno utjecati na ponašanje, a razlog tomu je pretpostavka da će unutar organizacije biti dostupna podrška i osposobljavanje za lakše rukovanje tehnologijom. S druge

strane, okruženja krajnjih potrošača razlikuju se po mnogim elementima, kao što su razlike u kvaliteti uređaja ili softvera kojim raspolažu i tomu slično. Dakle, u kontekstu korištenja tehnologije od strane krajnjeg potrošača olakšavajući uvjeti djelovat će poput percipirane bihevioralne kontrole u Teoriji planiranog ponašanja i utjecati i na namjeru (BI) i na stvarno korištenje (USE) (Venkatesh i sur., 2012; Ajzen, 1991).

Na primjer, ukoliko bi se promatralo korištenje alata generativne umjetne inteligencije, potrošači imaju različite razine pristupa brzom internetu, virtualnim instrukcijama na njihovom izvornom jeziku, kvaliteti računala koje koriste i slično. Potrošač koji ima povoljni skup olakšavajućih uvjeta vjerojatnije će imati veću namjeru korištenja tehnologije. S druge strane, stariji potrošači obično se suočavaju s većim poteškoćama u obradi novih ili složenijih informacija, što direktno utječe na njihovo učenje o korištenju novih tehnologija (Morris i sur., 2005). Stariji korisnici tehnologija pridaju veću važnost dostupnosti prikladne podrške nego mladi (Hall i Mansfield, 1975). Osim toga, muškarci su spremniji uložiti više napora u savladavanje raznih ograničenja i poteškoća u ostvarenju svojih ciljeva i manje se oslanjaju na olakšavajuće uvjete pri razmatranju korištenja tehnologije (Venkatesh i Morris, 2000; Venkatesh i sur., 2012). Nadalje, iskustvo također može moderirati odnos između olakšavajućih uvjeta i namjere ponašanja jer će iskusniji korisnici biti bolje upoznati s tehnologijom i imati bolju strukturu znanja zbog čega će se manje oslanjati na vanjsku podršku (Alba i Hutchinson, 1987; Venkatesh i sur., 2012). Imajući sve navedeno na umu, potkrijepljeno i prethodnim empirijskim dokazima (Morris i sur., 2005; Venkatesh i sur., 2003; Venkatesh i sur., 2012), zaključak je kako će starije žene posebno puno pažnje pridavati olakšavajućim uvjetima, a posebno one starije žene u ranim fazama korištenja određene tehnologije zbog čega Venkatesh i sur. (2012) u izvornom radu o UTAUT2 modelu navodi hipotezu koja glasi:

- *Dob, spol i iskustvo moderiraju utjecaj olakšavajućih elemenata na namjeru ponašanja, a posebno će utjecaj biti jak kod starijih žena u ranim fazama korištenja tehnologije.*

Hipoteza koja predviđa kako će dob, spol i iskustvo moderirati učinak olakšavajućih elemenata na namjeru ponašanja u izvornom radu o UTAUT2 modelu djelomično je podržana jer se pokazalo kako su samo spol i dob bili značajni moderatori, a iskustvo nije. Ipak, pretpostavka da će utjecaj biti posebno izražen kod starijih žena pokazala se točnom (Venkatesh i sur., 2012). Ova verzija hipoteze primjerenija je istraživanjima u neorganizacijskom okruženju nego li ona navedena kroz dvije podhipoteze u UTAUT-u koja se orijentira na olakšavajuće uvjete u organizacijskom okruženju.

Nadalje, Venkatesh i sur. (2012) navode kako se očekuje da će utjecaj konstrukta hedonističke motivacije na namjeru ponašanja biti moderiran dobi, spolom i iskustvom zbog razlika u inovativnosti potrošača, a što se djelomično može vezati i na konstrukt osobne inovativnosti (engl. *Personal innovativeness*) o čemu će biti više riječi u nastavku ovog rada. Naime, inovativnost je stupanj u kojem je pojedinac otvoren za nove ideje i samostalno donosi odluke o inovacijama, a traženje noviteta tendencija je pojedinca da traži nove informacije i podražaje (Midgley i Dowling, 1978; Hirschman, 1980). Takva inovativnost i traženje noviteta mogu doprinijeti hedonističkoj motivaciji za korištenje bilo kojeg proizvoda zbog čega će korisnici biti skloni posvetiti više pažnje novostima kod proizvoda i možda koristiti proizvod upravo zbog noviteta (Holbrook i Hirschman, 1982). Venkatesh i sur. (2012) na temelju prethodno navedenih elemenata i empirijskih dokaza (Lee i sur., 2010; Chau i Hui, 1998), da će hedonistička motivacija imati manje važnu ulogu u određivanju korištenja tehnologije s povećanjem iskustva jer s vremenom privlačnost noviteta opada pa se tehnologija krene više koristiti iz pragmatičnih razloga, ali i zbog činjenica da mlađi muškarci imaju tendenciju većeg traženja noviteta i inovativnosti. Shodno tomu, Venkatesh i sur. (2012) skrojili su hipotezu koja glasi:

- *Dob, spol i iskustvo moderirat će učinak hedonističke motivacije na namjeru ponašanja, tako da će učinak biti jači među mlađim muškarcima u ranim fazama iskustva s tehnologijom.*

U izvornom radu u kojem se UTAUT2 testira na prihvaćanju i korištenju mobilnog interneta, navedena hipoteza je podržana. Učinak je doista najjači među mlađim muškarcima u ranim fazama iskustva s tehnologijom (Venkatesh i sur., 2012).

Nadalje, Venkatesh i sur. (2012) navode kako se očekuje da će vrijednost cijene na namjeru ponašanja također biti moderirana dobi i spolom oslanjajući se na teorije o društvenim ulogama koje promatraju razlike između muškaraca i žena te mlađih i starijih pojedinaca i njihovoj percepciji o važnosti cijene proizvoda ili usluge (Bakan, 1966; Deaux i Lewis, 1984). Prethodno navedena istraživanja ukazuju na to da žene u pravilu obraćaju više pažnje na cijene proizvoda i usluge te da su svjesnije troškova nego li muškarci, no zanimljivo je da su muškarci skloniji dodijeliti veću vrijednost tehnologijama nego žene jer su općenito skloniji tehnologijama. Deaux i Lewis (1984) navode kako se navedeni stereotipi društvenih uloga dodatno pojačavanju starenjem tako da starije žene postaju sve osjetljivije na cijene zbog društvene uloge u kojoj su žene uglavnom te koje paze na kućni budžet. Imajući u vidu sve navedeno, Venkatesh i sur. (2012) navode hipotezu koja glasi:

- *Dob i spol će moderirati učinak vrijednosti cijene na namjeru ponašanja, tako da će učinak biti jači među ženama, posebno starijim ženama.*

U izvornom radu u kojem se UTAUT2 testira na prihvatanju i korištenju mobilnog interneta, navedena hipoteza je podržana. Učinak je doista najjači među starijim ženama (Venkatesh i sur., 2012).

Venkatesh i sur. (2012) navode kako je i treći novododani konstrukt, navika, također moderiran dobi, spolom i iskustvom. Postoje dva uzročna puta kroz koja navika konačno utječe na upotrebu, a oba se oslanjaju na obradu informacija i signala. U prvom, iskustvo uglavnom utječe na jačinu asocijacije između kontekstualnih signala i namjere ili ponašanja, a odnos između iskustva i navike formira se i jača kao rezultat ponovljenog ponašanja (Limayem i sur., 2007). Navika se promatra kao naučeni ishod koji se nakon relativno dugog razdoblja intenzivnog korištenja može pohraniti u dugoročno pamćenje i zamijeniti druge obrasce ponašanja (Lustig i sur., 2004; Limayem i Hirt, 2003). Stoga, navika ima jači učinak na namjeru i korištenje kod potrošača s više iskustva. Limayem i Hirt (2003) navode kako navika ima utjecaj i na namjeru ponašanja i korištenje tehnologije, ali se, kako vrijeme prolazi, utjecaj namjere ponašanja na stvarno ponašanje postupno smanjuje, a navika se postupno povećava. Slikovito rečeno, ukoliko učitelji svojim učenicima brzo usade naviku korištenja nove tehnologije, njihovi naponi u oblikovanju namjera učenika da koriste tehnologiju mogli bi s vremenom postati manje intenzivni. S druge strane, spol i dob odražavaju razlike ljudi u obradi informacija koje zauzvrat mogu utjecati na njihovu ovisnost o navici. Starije osobe uglavnom se oslanjaju na automatsku obradu informacija pri čemu ih navike sprječavaju u učenju novih stvari i jednom kada formiraju naviku teško ju nadvladaju da bi se prilagodili novim promjenama (Hasher i Zacks, 1979; Lustig i sur., 2004). Osim starenja, odnosno dobi, i razlike među spolovima igraju veliku ulogu pa su tako žene sklonije detaljno obrađivati informacije dok se muškarci, a posebno stariji muškarci, oslanjaju na heuristiku i sheme naučene iz iskustva (Darley i Smith, 1995). Imajući u vidu sve navedeno, Venkatesh i sur. (2012) skrojili su po jednu hipotezu za utjecaj navike na namjeru ponašanja i jednu za utjecaj navike na korištenje, a koje glase:

- *Dob, spol i iskustvo će moderirati učinak navike na namjeru ponašanja, tako da će učinak biti jači kod starijih muškaraca s visokim razinama iskustva s tehnologijom.*
- *Dob, spol i iskustvo će moderirati učinak navike na korištenje tehnologije, tako da će učinak biti jači kod starijih muškaraca s visokim razinama iskustva s tehnologijom.*

U izvornom radu u kojem se UTAUT2 testira na prihvatanju i korištenju mobilnog interneta, obje navedene podhipoteze su podržane (Venkatesh i sur., 2012). Doista, učinak se pokazao jačim među starijim muškarcima s više iskustva u korištenju mobilnog interneta i po pitanju namjere ponašanja, ali i stvarnog ponašanja kod korištenja tehnologije.

Nadalje, Venkatesh i sur. (2012) navode kako u UTAUT2 modelu postoji i utjecaj namjere ponašanja na ponašanje korisnika tehnologije, a koji je moderiran iskustvom. Spomenuti autori navode kako s povećanjem iskustva potrošači imaju više prilika za pojačati svoju naviku, a samim time rutinsko ponašanje postaje automatizirano (Venkatesh i sur., 2012; Jasperson i sur., 2005). Sumirajući navedeno, Venkatesh i sur. (2012) navode hipotezu koja glasi:

- *Iskustvo će moderirati učinak namjere ponašanja na korisničko ponašanje kod korištenja tehnologije, tako da će učinak biti jači za potrošače s manje iskustva.*

U izvornom radu u kojem se UTAUT2 testira na prihvatanju i korištenju mobilnog interneta, navedena hipoteza je podržana. Učinak doista opada kako se povećava iskustvo potrošača (Venkatesh i sur., 2012).

Venkatesh i sur. (2012) navode kako je najveći doprinos UTAUT2 modela u odnosu na izvorni UTAUT što je on sada prilagođen za istraživanja u kontekstu prihvatanja i korištenja tehnologije od strane krajnjih potrošača, odnosno teorijski je proširena primjenjivost modela sa organizacijskog konteksta na potrošački kontekst. UTAUT i UTAUT2 samo su dobar teorijski okvir za istraživanja prihvatanja i korištenja tehnologije, ali oni su pogodni i za proširivanja relevantnim konstruktima i čimbenicima koji će povećati primjenjivost UTAUT-a u raznim kontekstima. Navedeni autori navode: „buduća istraživanja mogu se graditi na našem istraživanju testiranjem UTAUT2 modela u različitim zemljama, različitim dobima i različitim tehnologijama“ (Venkatesh i sur., 2012). Nadalje, navedeni autori napominju kako je pitanje varijance zajedničke metode (CMV) glavna metodološka briga povezana s ovim istraživanjem te da buduća istraživanja mogu koristiti i neki rigorozniji dizajn kako bi smanjili mjernu i metodičku pristranost.

4. PREGLED LITERATURE

Venkatesh i sur. (2012) navode kako je u izvornom radu u kojem je predstavljen UTAUT2, a u kojem se ispituje prihvatanje i korištenje mobilnog interneta, većina hipoteza podržana te kako UTAUT2 sa samo izravnim učincima objašnjava 44 % varijance u namjeri ponašanja, a UTAUT2 koji uključuje moderacijske efekte objašnjava 74 % varijance u namjeri ponašanja. Kada se promatra objašnjenje korištenja tehnologije, model UTAUT2 sa samo izravnim učincima objašnjava 35 % varijance dok s uključenim interakcijskim pojmovima objašnjava 52 % varijance. Sve to predstavlja znatno unaprjeđenje u odnosu na izvorni UTAUT.

Rad pod nazivom „Petrošačko prihvatanje i korištenje informacijske tehnologije: proširenje ujedinjene teorije prihvatanja i korištenja tehnologije“ u kojem su Venkatesh i sur. (2012) predstavili UTAUT2 model danas (u svibnju 2025. godine) ima citiranost od gotovo 20 tisuća spominjanja na Google znalcu. Navedeni rad istraživačima iz područja informacijskih sustava bitan je u smislu teorijske perspektive za razumijevanje pitanja povezanih s prihvatanjem tehnologije u različitim kontekstima, bilo samostalno ili u kombinaciji s drugim teorijama i dodatnim vanjskim varijablama (Tamilmani i sur., 2021).

U ranijim studijama u kojima su rađeni detaljni pregledi razvoja ovog modela, uočeno je kako 77 % znanstvenih i stručnih članaka citira UTAUT2 u opće svrhe, a preostalih 23 % to čini u kombinaciji s vanjskim teorijama i rijetkim uključivanjem moderatora (Tamilmani i sur., 2017; Tamilmani i sur., 2021). Tamilmani i sur. (2021) navode kako su istraživači u svojim nastojanjima da razumiju individualno prihvatanje i korištenje tehnologije primjenili, integrirali i proširili UTAUT2 u raznim kontekstima, a što se može podijeliti u šest kategorija:

- različite vrste korisnika
- različite vrste organizacija
- različite vrste tehnologija
- različite vrste zadataka
- različita vremena
- različite lokacije.

Marikyan i Papagiannidis (2023) navode kako je UTAUT tijekom godina pokazao široku primjenu čime je pojačana generalizabilnost teorije jer su kroz vrijeme znanstvenici proširivali model kako bi ga prilagodili kontekstu i poboljšali njegovu prediktivnu snagu. UTAUT2 izvorno nije ni dizajniran s posebnim naglaskom na određenu novu tehnologiju ili geografsku

lokaciju. Cilj modela bio je predstaviti sveobuhvatan okvir za ispitivanje prihvaćanja i korištenja tehnologije, a navedena proširenja u konceptualnom dizajnu modela koriste se za povećanje preciznosti u objašnjavanju ponašanja korisnika ovisno o kontekstu onoga što se ispituje (Marikyan i Papagiannidis, 2023; Venkatesh i sur., 2012).

Tablica 9. Pregled literature

Redni broj	Izvor	Model	Objekt	Faktori	Modifikacije
1.	Ain i sur. (2016)	UTAUT2	Sustavi za upravljanje učenjem (LMS) (Malezija)	PE, EE, SI, FC, HM, LV , HT BI, USE	Umjesto PV - <i>Price Value</i> koristi LV - <i>Learning Value</i> (vrijednost učenja).
2.	Al-Emran (2023)	ICAAM	Chatbotovi pogonjeni umjetnom inteligencijom (Malezija)	PE, EE, SI, FC, HM, HT, PT , PS , PSUS CHATBOT USE	Nije UTAUT2 nego srodan model. Nema <i>Price Value</i> i nema <i>Behavioral intention</i> , ali su dodani PT - <i>Perceived Threat</i> (Percipirana prijetnja), PS - <i>Perceived Severity</i> (Percipirana ozbiljnost) i PSUS - <i>Perceived Susceptibility</i> (Osjetljivost).
3.	Ameri i sur. (2020)	UTAUT2	Mobilne aplikacije za obrazovanje u farmaciji (LabSafety) (Iran)	PE, EE, SI, FC, HT BI > USE	Nema faktora PV - <i>Price Value</i> i HM - <i>Hedonic Motivation</i> .
4.	Baabdullah (2018)	UTAUT2	Mobilne igre na društvenim mrežama (Saudijska Arabija)	PE, EE, SI, FC, HM, PV, TR BI	Rad koristi UTAUT2 model iz kojeg je izostavljen konstrukt HT, te zavisna varijabla USE. Također, rad ne koristi ni moderatore. U model je dodan konstrukt <i>Trust</i> – TR (Povjerenje).
5.	Baudier i sur. (2020)	UTAUT2	Smart Home koncept (Francuska i Južna Koreja)	PE, EE, SI, FC, HM, PV, HT, PI , SS , H , CC , S INTENTION TO USE (BI)	Uz standardne faktore UTAUT2 teorije dodani su SS - <i>Safety Security</i> (Sigurnost), H - <i>Health</i> (Zdravlje), CC - <i>Convenience Comfort</i> (Praktična udobnost), S - <i>Sustainability</i> (Održivost) kao dio Smart Home dimenzije istraživanja te PI - <i>Personal Innovativeness</i> (Osobna inovativnost) kao dio IT dimenzije.
6.	Baptista & Oliveira (2015)	UTAUT2	Mobilno bankarstvo (Mozambik)	PE, EE, SI, FC, HM, PV, HT BI > USE	Rad koristi UTAUT2 model koji je proširen dodanim moderatorima s kulturološkim kontekstom (individualizam/kolektivizam, neizvjesnost/izbjegavanje, dugoročnost/kratkoročnost, muževnost/ženstvenost).

7.	Bellet & Banet (2023)	UTAUT4-AV	Autonomna vozila (Francuska)	PU, PE, EE, HM, SI, PS, PV, ANX, PV2*, SAT INTENTION TO USE (BI)	Na standardne faktore UTAUT2 teorije dodani su PU - <i>Perceived usefulness</i> (Percipirana korisnost), PS - <i>Perceived Safety</i> (percipitana sigurnost), ANX - <i>Anxiety</i> (anksioznost), SAT - <i>Satisfaction</i> (zadovoljstvo) te drugi faktor cijene koji uspoređuje sadašnji sustav sa autonomnim sustavom. Dodani su i mnogobrojni medijatori.
8.	Biloš & Budimir (2024)	UTAUT2	ChatGPT (Hrvatska)	PE, EE, SI, FC, HM, PV, HT, PI BI	Na standardne faktore UTAUT2 teorije dodan je prediktor PI - <i>Personal Innovativeness</i> (Osobna inovativnost), a izbačeni su moderatori.
9.	Budhathoki i sur. (2024)	UTAUT	ChatGPT (Nepal i UK)	PE, EE, SI, FC, ANX BI > USE	Rad koristi UTAUT model, navodi 4 standardna prediktora + anksioznost – ANX (<i>Anxiety</i>). Rad ne koristi moderatore.
10.	Cabrera-Sanchez i sur. (2021)	UTAUT2	Primjena umjetne inteligencije (Španjolska)	PE, EE, SI, PV, HT, FC, TF, CT BI > USE	Izbačeni su moderatori dob, spol i iskustvo, a dodani su TF - <i>Technology Fear</i> (strah od tehnologije) i CT - <i>Consumer Trust</i> (Potrošačko povjerenje) kao prediktori, ali i moderatori.
11.	Cao i sur. (2023)	TAM	Primjena umjetne inteligencije (Kina)	<i>Perceived Usefulness, Perceived Ease of Use, Knowledge or Skills in AI, Attention to AI news, Attitude toward Using AI, Behavioral Intention and Actual System Use</i>	Nije korišten UTAUT2 već modificirani TAM koji koristi: <i>Perceived Ease of Use</i> koji je srodan konstrukt EE, <i>Perceived Usefulness</i> koji je srodan konstrukt PE, ali ovaj rad koristi i druge konstrukte kao što su <i>Knowledge or Skills in AI, Attention to AI news, Attitude toward Using AI</i> (Stav prema korištenju AI) te namjera korištenja (BI) i stvarno korištenje (USE).
12.	Chatterjee & Bhattacharjee (2020)	UTAUT	Primjena umjetne inteligencije (Indija)	PR, PE, EE, FC, ATT > BI > USE	U radu se koristi UTAUT, no umjesto konstrukta SI (Društvenog utjecaja) korišten je konstrukt PR – <i>Perceived Risk</i> (Percipirani rizik). Osim toga u rad je

					dodana i zavisna varijabla ATT (engl. <i>Attitude</i> - Stav prema tehnologiji)
13.	Chopdar i sur. (2018)	UTAUT2	Aplikacije za m-trgovinu (Indija i SAD)	PE, EE, SI, FC, HM, PV, HT, PR, SR BI > USE	Rad koristi UTAUT2 model uz ekstenzije Rizik po pitanju privatnosti – PR (<i>Privacy risk</i>) i Rizik po pitanju sigurnosti – SR (<i>Security risk</i>). HT, PR i SR, prediktori su za BI i za USE, ostali prediktori samo utječu na BI.
14.	Choudhary i sur. (2024)	BRT	Virtualni glasovni asistenti pogonjeni umjetnom inteligencijom (Indija)	PE, EE, SI, FC, PV, HM, PR, VB, IB, VOC, RF, RA ATT > BI	Rad koristi modificirani BRT model (<i>Behavioral Reasoning Theory</i>) koji je proširen elementima iz UTAUT2 modela. Korišteni su svi prediktori UTAUT2 modela osim vrijednosti cijene, ali osim toga u modelu se nalaze: Otvorenost prema promjenama - VOC (<i>Value of openness to change</i>), Razlozi za prihvatanje (RF), Razlozi protiv prihvatanja (RA), Percipirani rizik (PR), <i>Image barrier</i> (IB – prepreka zbog imidža) i <i>Value Barrier</i> (VB – vrijednosna prepreka).
15.	Chow i sur. (2023)	-	Primjena umjetne inteligencije unutar FinTech sustava (Hong Kong)	PE, EE, <i>Cybersecurity, Anthropomorphism, Online Banking Frequency, Compensatory Level of Acceptance to AI services</i>	Nije korišten niti jedan model, već je konceptualni model sastavljen na temelju više teorija koje proučavaju prihvatanje tehnologije.
16.	Chu i sur. (2022)	UTAUT2	Inteligentna dizala (Tajvan)	PE, EE, SI, FC, HT, EC, AIO, PQ ATT > BI	Radi koristi modificirani UTAUT2 model u kojem nema zavisne varijable USE, izbačeni su prediktori PV i HM, ali su dodani Osviještenost o okolišu - EC (<i>Environmental consciousness</i>), Optimizam u vezi umjetne inteligencije - AIO (<i>Artificial intelligence optimism</i>) te percipirana kvaliteta – PQ (<i>Perceived quality</i>).
17.	Cintron (2022)	UTAUT2	Primjena umjetne inteligencije (SAD)	PE, EE, SI, FC, PV, HT BI	Rad koristi modificirani UTAUT2 u kojem nema zavisne varijable USE te prediktora HM.

18.	Cortez i sur. (2024)	UTAUT2 + SDT	Komunikacijska umjetna inteligencija u obrazovanju (Filipini)	PE, EE, FC, HM, PV, HT, PA, PC, PR BI > USE	Rad koristi UTAUT2 model u kombinaciji sa SDT (<i>Self-determination theory</i>). Iz modela je izostavljen faktor SI, dodani su Percipirana autonomija – PA (<i>Perceived autonomy</i>), Percipirana kompetentnost – PC (<i>Perceived competency</i>) i Percipirana povezanost – PR (<i>Perceived relatedness</i>). Rad ne koristi moderatore.
19.	Faqid & Jaradat (2021)	TTF + UTAUT2	Proširena stvarnost (AR) (Jordan)	PE, EE, SI, FC, HM, PV, TECH-C, TASK-C, TTF BI	U radu je korišten model koji spaja TTF (<i>Technology Task Fit model</i>) i UTAUT2 (bez faktora USE, HT i moderatora). TTF sastoji se od faktora TECH-C (<i>Technology characteristics</i>), TASK-C (<i>Task characteristics</i>) i TTF (<i>Task-Technology Fit</i>).
20.	Foroughi i sur. (2024)	UTAUT2	ChatGPT (Malezija)	PE, EE, SI, FC, HM, HT LV, PI, IA BI	Rad koristi UTAUT2 model, umjesto konstrukta <i>Price Value</i> dodan je konstrukt <i>Learning Value</i> – LV (Vrijednost učenja), a dodani su Osobna inovativnost - PI (<i>Personal Innovativeness</i>) i Točnost informacija - IA (<i>Information Accuracy</i>). U modelu se PI i IA koriste kao moderatori.
21.	Gansser & Reich (2021)	UTAUT2	Primjena umjetne inteligencije (Njemačka)	PE, EE, SI, PV, HM, HT, CC, S, SS, PI BI > USE	Rad koristi prošireni UTAUT2 model. Nema prediktora FC, ali su dodani Praktična udobnost – CC (<i>Convenience comfort</i>), Održivost – S (<i>Sustainability</i>), Sigurnost – SS (<i>Safety security</i>) i Osobna inovativnost – PI (<i>Personal innovativeness</i>).
22.	Garcia de Blanes Sebastian i sur. (2022)	UTAUT2	Virtualni glasovni asistenti pogonjeni umjetnom inteligencijom (Španjolska)	PE, EE, SI, FC, PV, HM, HT, PPR, PT, PI BI	Rad koristi UTAUT2 model proširen s tri dodatna prediktora: Percipirani rizik po pitanju privatnosti – PPT (<i>Perceived privacy risk</i>), Percipirano povjerenje – PT (<i>Perceived trust</i>) i Osobna inovativnost – PI (<i>Personal innovativeness</i>). Također, u radu nema faktora USE niti moderatora.

23.	Goyal i sur. (2023)	UTAUT2	ChatGPT (Indija)	PE, EE, SI, FC, HM, HT PC > ATT > BI	Rad koristi modificirani UTAUT2 model u kojem nema prediktora PV, ali preostalih šest faktora prediktori su Percipirane vjerodostojnosti – PC (<i>Perceived credibility</i>), a što utječe na Stav prema korištenju – ATT (<i>Attitude</i>), a što u konačnici utječe na namjeru korištenja – BI.
24.	Gursoy i sur. (2019)	AIDUA	Uređaji pogonjeni umjetnom inteligencijom (SAD)	PE, EE, SI, HM, ANTH, EM, WTA, OTU	U radu nije korišten UTAUT2 model već AIDUA model. Iz UTAUT2 modela koristi prediktore PE, EE, SI, HM, a dodani su Antropomorfizam (ANTH – <i>Anthropomorphism</i>), Emocije (EM – <i>Emotion</i>), Volja za prihvatanjem umjetne inteligencije (WTA - <i>Willingness to accept the use of AI Devices</i>) i Odbijanje korištenja umjetne inteligencije (OTU - <i>Objection to Use of AI Devices</i>).
25.	Gotheresen i sur. (2023)	UTAUT2	Tehnologija za pametnu kuću (Norveška)	PE, EE, SI, HM, PV, FC, S/P, EM BI	Rad koristi UTAUT2 model iz kojeg je izostavljen konstrukt HT te faktor USE. Također, izostavljeni su i moderatori, ali su dodani prediktori Sigurnost/Privatnost – S/R (<i>Security/Privacy</i>) i Upravljanje energijom – EM (<i>Energy management</i>)
26.	Grassini i sur. (2024)	UTAUT2	ChatGPT (Norveška)	PE, EE, SI, FC, HM, HT BI > USE	Rad koristi UTAUT2 model iz kojeg je izostavljen konstrukt PV. Model koristi svih 6 preostalih prediktora te standardne moderatore faktore BI i USE.
27.	Habibi i sur. (2023)	UTAUT2	ChatGPT (Indonezija)	PE, EE, SI, HM, FC, HT BI > USE	Rad koristi UTAUT2 model, izostavljen je prediktor PV te moderatori nisu korišteni u modelu.
28.	Hassan i sur. (2023)	UTAUT2	FinTech usluge (Bangladeš)	PE, EE, SI, FC, PV, PC <i>Acceptance of MFS</i>	U radu je korišten modificirani UTAUT2 model u kojem su izostavljeni HT i HM te moderatori, dodan je faktor percipirane vjerodostojnosti – PC (<i>Perceived credibility</i>), a umjesto BI i USE koristi se faktore prihvatanja što je nešto između BI i USE, ali bliži namjeri korištenja.

29.	Huang i sur. (2024)	UTAUT2	Umjetna inteligencija u medicini (Singapur)	PE, EE, SI, FC, HM BI	Rad koristi UTAUT2 model iz kojeg su izostavljeni konstrukti PV i HT, a ne koristi se niti faktor USE niti moderatori.
30.	Kasparova (2022)	UTAUT2	Alati za poslovnu inteligenciju (Češka)	PE, EE, SI, FC, HT BI > USE	Rad koristi modificirani UTAUT2 model iz kojeg su izostavljeni PV i HM, a svi preostali faktori osim PE i EE su ujedno prediktori i za BI i za USE. Također, u radu su izostavljeni i moderatori.
31.	Kaya i sur. (2024)	GAAIS	Primjena umjetne inteligencije (Turska)		U radu se koristi GAAIS model (<i>General Attitudes toward Artificial Intelligence Scale</i>) koji ne koristi faktore koji su srodni onima u UTAUT2 modelu, ali istražuju sklonost ka korištenju umjetne inteligencije.
32.	Khan i sur. (2022)	UTAUT3	Primjena sustava za mobilno učenje	PE, EE, SI, FC, HM, HT, PV, PI BI	Autor korišteni model naziva UTAUT3, riječ je o UTAUT2 modelu proširenim za faktor Osobne inovativnosti – PI (<i>Personal innovativeness</i>). Također, proširen je i popis moderatora koji utječu na namjeru korištenja, prilagođeni za istraživanja na studentskoj populaciji.
33.	Kilani i sur. (2023)	UTAUT2	Korištenje digitalnih novčanika (Jordan)	PE, EE, FC, HM, PV, HT, TR UI (BI) > CUB (USE)	Rad koristi UTAUT2 model iz kojeg je izostavljen prediktor društvenog utjecaja (SI), ali je dodan prediktor Povjerenja – TR (<i>Trust</i>). Također, iz modela su izostavljeni i moderatori.
34.	Korkmaz i sur. (2022)	UTAUT2	Sustav autonomnih vozila u javnom prijevozu (Turska)	PE, EE, SI, FC, HM, PV, HT, PU, TR, S, PR BI	Rad koristi UTAUT2 model sa svih 7 standardnih prediktora za BI, ali je izostavljen faktor USE. U radu se koriste moderatori dob, spol i učestalost. Osim navedenog, modelu su pridruženi konstrukti Percipirane korisnosti – PU (<i>Perceived usefulness</i>), Povjerenje – TR (<i>Trust</i>), Sigurnost – S (<i>Safety</i>), Percipirani rizik – PR (<i>Perceived risk</i>).

35.	Lai i sur. (2024)	UTAUT	ChatGPT (Hong Kong)	PE, EE, SI, MO, TR, PR BI	Rad ne koristi UTAUT2 model već prvi UTAUT. Iz modela su izostavljeni moderatori te faktor USE. Modelu si pridodani konstrukt Moralna dužnost – MO (<i>Moral obligation</i>), Povjerenje – TR (<i>Trust</i>) i Percipirani rizik – PR (<i>Perceived risk</i>).
36.	Lavidas i sur. (2024)	UTAUT2	ChatGPT (Grčka)	PE, EE, SI, FC, HM, PV, HT BI > USE	Rad koristi standardni UTAUT2 model uključujući svih 7 prediktora za BI, koristi se i faktor USE te tri standardna moderatora (dob, spol i iskustvo).
37.	Lee i sur. (2024)	UTAUT	ChatGPT (SAD)	PE, EE, SI, FC, RRP, NE ATT > BI	Rad ne koristi UTAUT2 model već prvi UTAUT te sva četiri standardna prediktora iz navedenog modela uz kojeg su nadodani Percepcija relativnog rizika – RRP (<i>Relative risk perception</i>) i Negativna emocija – NE (<i>Negative emotion</i>) te Stav – ATT (<i>Attitude</i>) kao faktor na koji utječu navedeni prediktori. Rad koristi moderatore – dob, spol, obrazovanje, etnička skupina i prihodi.
38.	Limayem & Hirt (2003)	TPB + Triandisov model	Interaktivna internetska komunikacijska tehnologija (Hong Kong)	<i>Habit, Social factors,</i> <i>Facilitating conditions,</i> <i>Perceived consequences,</i> <i>Affect</i> <i>Intentions > Usage behavior</i>	Rad koristi spoj Planirane teorije ponašanja (TPB) i Triandisovog modela i predstavlja jedan o modela koji je pomogao u stvaranju UTAUT2 modela.
39.	Macedo (2017)	UTAUT2	Informacijsko- komunikacijska tehnologija (Portugal)	PE, EE, SI, FC, HM, PV, HT BI > USE	Rad koristi UTAUT2 model sa svih sedam standardnih prediktora za BI te i faktor USE. U radu se umjesto moderatore koriste kontrolne varijable – dob, spol, iskustvo i obrazovanje.
40.	Maican i sur. (2023)	UTAUT2	Primjena generativne umjetne inteligencije za	PE, EE, PCV , HM, SI, FC, HT	Rad koristi UTAUT2 model iz kojeg su izostavljeni prediktor PV, model nema faktor USE, ali je dodan faktor Percipirana vrijednost za korisnika – PCV (<i>Perceived customer value</i>). Nadalje, rad koristi dodatne moderatore

			slike (Rumunjska)	BI	kao što su spol, kreativnost i razina znanja engleskog jezika te medijatore na BI koji se dotiču različitih primjena generativne AI za slike (inženjerstvo, znanost, umjetnost).
41.	Maruping i sur. (2017)	UTAUT	Informacijska tehnologija (SAD)	PE, EE, SI, FC BI > USE, BE > USE	Rad ne koristi UTAUT2 već prvi UTAUT uz dodatak konstrukta Očekivanog ponašanja – BE (<i>Behavioral Expectation</i>). U radu su korišteni standardni moderatori za UTAUT – spol, dob, iskustvo i dobrovoljnost korištenja.
42.	Meet i sur. (2022)	UTAUT2	Sustavi za online tečajeve (Indija)	PE, EE, SI, FC, HM, PV, HT, LC, TI BI	Rad koristi prošireni UTAUT2 model, u kojem su uz standardnih sedam faktora dodani i jezična kompetencija – LC (<i>Language competency</i>) i utjecaj nastavnika – TI (<i>Teacher influence</i>).
43.	Mehedi Hasan Emon i sur. (2023)	UTAUT2	ChatGPT (Bangladeš)	PE, EE, SI, FC, HM, TR, ATT-AI BI > USE	Rad koristi UTAUT2 model iz kojeg su izostavljeni konstrukti PV i HT kao prediktori. Rad ne koristi moderatore. Modelu su pridodani konstrukti Povjerenje – TR (<i>Trust</i>) i Stav prema umjetnoj inteligenciji – ATT-AI (<i>Attitude toward Artificial Intelligence</i>).
44.	Menon & Shilpa (2023)	UTAUT	ChatGPT (Indija)	PE, EE, SI, FC, PI, PC, PHT	Rad koristi UTAUT model, ali je specifičan po tome što se provodi kao kvalitativno istraživanje. Osim standardna četiri konstrukta dodani su Percipirana interaktivnost - PI (<i>Perceived interactivity</i>), Zabrinutost za privatnost – PC (<i>Privacy concerns</i>) i Percipirani ljudski utjecaj (<i>Perceived human touch</i>).
45.	Merhi i sur. (2019)	UTAUT2	Mobilno bankarstvo (Libanon i UK)	PE, EE, SI, HM, PV, HT, TR, PP, PS BI	U radu je korišten prošireni UTAUT2 model iz kojeg je izostavljen konstrukt FC, ali su dodani Povjerenje – TR (<i>Trust</i>), Percipirana privatnost – PP (<i>Perceived Privacy</i>) i Percipirana sigurnost – PS (<i>Perceived Security</i>). U radu nema moderatora niti je korišten faktor USE.

46.	Miladinović & Hong (2016)	UTAUT2	Aplikacije za m-trgovinu za modu (Švedska)	PE, EE, SI, FC, HM, PV, HT, TR BI	U radu je korišten UTAUT2 model kojem je dodan konstrukt Povjerenje – TR (<i>Trust</i>), a izostavljen je faktor USE. Također, izostavljeni su i moderatori.
47.	Mohd Rahim i sur. (2022)	UTAUT2	Chatbot (Malezija)	PE, EE, SI, FC, HT, HM, PT, D, E, I BI > USE	U radu je korišten UTAUT2 model proširen konstruktima Percipirano povjerenje – PT (<i>Perceived trust</i>), Dizajn – D (<i>Design</i>), Etika – E (<i>Ethics</i>) i Interaktivnost – I (<i>Interactivity</i>). D, E i I prediktori su za PT. Iz rada su izostavljeni moderatori i PV.
48.	Nikolopoulou i sur. (2020)	UTAUT2	Mobilni telefoni (Grčka)	PE, EE, SI, FC, HM, PV, HT BI > USE	U radu je korišten standardni UTAUT2 model, identičan kao u izvornom radu.
49.	Nordhoff i sur. (2020)	UTAUT2	Autonomni automobili (8 europskih država)	PE, EE, SI, FC, HM BI	U radu je korišten UTAUT2 model iz kojeg su izostavljeni konstrukti HT i PV, ali je model modificiran tako da svi faktori utječu na sve ostale. Nadalje, u radu se koriste standardni moderatori za UTAUT2, ali se ne koristi faktor USE.
50.	Oliveira i sur. (2016)	UTAUT2 + DOI	Mobilno plaćanje (Portugal)	PE, EE, SI, FC, HM, PV, C, I, PTS BI-A > BI-R	Rad koristi UTAUT2 model u kombinaciji sa DOI modelom (<i>Diffusion of innovation</i>). Koristi se 6 prediktora iz UTAUT2 modela, svi osim HT. Modelu su dodani Kompatibilnost – C (<i>Compatibility</i>), Inovativnost – I (<i>Innovativeness</i>), Percipirana sigurnost tehnologije – PTS (<i>Perceived technology security</i>). Nadalje, rad koristi dvije razine namjere – Namjera prihvatanja tehnologije i Namjera preporuke tehnologije. Rad ne koristi moderatore.
51.	Salifu i sur. (2024)	UTAUT2	ChatGPT (Gana)	PE, EE, SI, FC, HT, HM, PT, D, E, I	U radu je korišten UTAUT2 model proširen konstruktima Percipirano povjerenje – PT (<i>Perceived trust</i>), Dizajn – D (<i>Design</i>), Etika – E (<i>Ethics</i>) i Interaktivnost – I

				BI > USE	(Interactivity). D, E i I prediktori su za PT. Iz rada su izostavljeni moderatori i PV.
52.	Sallam i sur. (2023)	Prošireni TAM	ChatGPT	PU, PR, PEOU, B/C, PR, ATT, ANX	U radu nije korišten UTAUT2 model već srodni TAM koji je proširen dodanim konstruktima. Srodna je tema jer istražuje prednosti korištenja ChatGPT-a kod studenata medicine.
53.	Schmitz i sur. (2022)	UTAUT2	Telemedicina i virtualni liječnički pregledi (Njemačka i SAD)	PE, EE, SI, FC, HT, HM, PS, PPA USAGE INTENTION (BI)	U radu je korišten UTAUT2 model iz kojeg je izostavljen konstrukt PV, ali su dodani Percipirana sigurnost – PS (<i>Perceived security</i>) i Percipirana prednost proizvoda – PPA (<i>Perceived product advantage</i>). Od moderatora su korišteni dob i spol.
54.	Setiyani i sur. (2023)	UTAUT2	E-trgovina (Indonezija)	PE, EE, SI, SC (FC), HM, PV, HT BI	Rad koristi UTAUT2 model iz kojeg su izostavljeni moderatori te faktor USE. Unutar modela konstrukt FC naziva se <i>Supporting conditions</i> umjesto <i>Facilitating conditions</i> kako je nazvan u izvornom radu.
55.	Shaw & Sergueeva (2019)	UTAUT2	M-trgovina (Kanada)	PE, EE, SI, FC, HM, HT, PV, PI, PPC, PPR, PTR, PPP INTENTION TO USE (BI)	Rad koristi UTAUT2 model iz kojeg su izostavljeni njegovi standardni moderatori. U radu se koristi samo moderator Osobne inovativnosti (PI) koji moderira utjecaj PE na Percipiranu vrijednost - PV (<i>Perceived Value</i>) i Percipiranu zabrinutost po pitanju privatnosti – PPC (<i>Perceived privacy concerns</i>). Na PPC utječu Percipirani rizik po pitanju privatnosti – PPR (<i>Perceived Privacy Risk</i>), Percipirani transakcijski rizik – PTR (<i>Perceived transaction risk</i>) i Percipirana zaštita privatnosti - PPP (<i>Perceived Privacy Protection</i>).
56.	Shi i sur. (2022)	UTAUT2	Internet stvari u poljoprivredi (Bangladeš)	PE, EE, SI, FC, HM, PV, TR, GS, PI,	U radu je korišten modificirani UTAUT2 bez moderatora. Umjesto faktora BI i USE koriste se Spremnost na prihvatanje – WTA (<i>Willingness to adopt</i>) i Spremnost na plaćanje – WTP (<i>Willingness to pay</i>). Osim toga dodani su konstrukti Osobne inovativnosti – PI (<i>Personal</i>

				WTA > WTP	<i>innovativeness</i>), Povjerenje – TR (<i>Trust</i>) te Vladina pomoć – GS (<i>Government support</i>).
57.	Sinaga i sur. (2024)	UTAUT2	ChatGPT (Indonezija)	PE, EE, SI, FC, HM, PV, HT, C BI > USE	Rad koristi UTAUT2 model koji uključuje svih sedam standardnih prediktora uz dodatak konstrukta Znatizeljnost – C (<i>Curiosity</i>) za BI, koristi se faktor USE te moderatori (dob, spol i iskustvo).
58.	Sobaih i sur. (2024)	UTAUT	ChatGPT (Saudijska Arabija)	PE, EE, SI, FC BI > USE	Rad ne koristi UTAUT2 model već prvi UTAUT. Koriste se sva četiri standardna prediktora te faktori BI i USE. Jedina modifikacije je što su iz modela izbačeni moderatori.
59.	Strzelecki (2023)	UTAUT2	ChatGPT (Poljska)	PE, EE, SI, FC, HM, HT, PI BI > USE	Rad koristi UTAUT2 model iz kojeg je izostavljen konstrukt PV, ali je dodan konstrukt Osobne inovativnosti – PI (<i>Personal innovativeness</i>). Korišteni moderatori su godina na studiju i spol.
60.	Sugumar & Chandra (2021)	UTAUT2 + BDI	Chatbot za financijske usluge (Indija, SAD i Singapur)	PE, EE, SI, HM, HT, ANTH, LA, SP ADOPTION INTENTION (BI)	Rad koristi spoj UTAUT2 modela i BDI modela. Iz UTAUT2 modela izostavljeni su PV i FC. Od BDI modela dodani su Antropomorfizam – AM (<i>Anthropomorphism</i>), Dopadljivost – LA (<i>Likeability</i>) i Društvena prisutnost – SP (<i>Social presence</i>). Kontrolne varijable su dob, spol, iskustvo sa <i>chatbotovima</i> za financijske usluge i iskustvo s bilo kakvim <i>chatbotovima</i> .
61.	Sun i sur. (2023)	UTAUT2* (potpuno modificiran)	Brza dostava pića (Kina)	SI, HT, PI, PR, K USAGE (USE)	Autori model nazivaju UTAUT2 iako je potpuno izmijenjen. Od standardnih UTAUT2 faktora preostali su samo SI i HT, a dodani su Inovativnost što je srodno kao Osobna inovativnost – PI, dodan je konstrukt Percipiranog rizika – PR (<i>Perceived Risk</i>) te je dodan konstrukt Znanja – K (<i>Knowledge</i>).

62.	Tantra & Ariyanti (2017)	UTAUT2	Mobilna aplikacija za akademski informacijski sustav (Indonezija)	PE, EE, SI, FC, HM, HT, CQ BI	Rad koristi UTAUT2 model iz kojeg su izostavljeni moderatori i faktor USE. Umjesto konstrukta PV dodan je konstrukt Kvaliteta sadržaja – CQ (<i>Content quality</i>)
63.	Tarhini i sur. (2024)	UTAUT2 + ISS	Servisi na mobilno učenje (Oman)	PE, EE, SI, FC, HM, PV, HT, P, SQ, IQ BI	Rad koristi kombinaciju UTAUT2 modela i ISS (<i>Information systems success</i>) modela. U radu se koristi svih sedam standardnih UTAUT2 prediktora, a njima su pridodani Privatnost – P (<i>Privacy</i>), Kvaliteta sustava – SQ (<i>System quality</i>) i Kvaliteta informacija – IQ (<i>Information quality</i>). U radu se ne koriste moderatori niti faktor USE.
64.	Tseng i sur. (2022)	UTAUT2	Sustavi za online tečajeve (Tajvan)	PE, EE, SI, FC, PV, HM BI > USE	Rad koristi UTAUT2 model iz kojeg je izostavljen konstrukt HT. Također, iz rada su izostavljeni i moderatori.
65.	Thusi & Maduku. (2020)	UTAUT2	Aplikacije za mobilno bankarstvo (JAR)	PE, EE, SI, FC, HM, PV, HT, IBT, PR, FR, PFR, PR, PSY-R, DR BI > USE	Rad koristi značajnije modificiran i proširen UTAUT2 model koji osim sedam standardnih prediktora koristi konstrukte Povjerenje u instituciju – IBT (<i>Institution-based trust</i>), Percipirani rizik – PR (<i>Perceived risk</i>) kao faktor na koji utječu Financijski rizik – FR (<i>Financial risk</i>), Rizik izvedbe – PFR (<i>Performance risk</i>), Rizik po pitanju privatnosti – PR (<i>Privacy risk</i>), Psihološki rizik – PSY-R (<i>Psychological risk</i>) i Rizik uređaja – DR (<i>Device risk</i>)
66.	Vimalkumar i sur. (2021)	UTAUT2	Virtualni glasovni asistenti (Indija)	PE, EE, SI, FC, HM, PV, PT, PPR, PPC BI > USE	Rad koristi modificirani UTAUT2 model iz kojeg je izostavljen faktor HT, ali su dodani Percipirano povjerenje – PT (<i>Perceived Trust</i>), Percipirani rizik po pitanju privatnosti – PPR (<i>Perceived Privacy Risk</i>) i Zabrinutost po pitanju privatnosti - PPC (<i>Perceived Privacy Concerns</i>)
67.	Vinerean i sur. (2022)	UTAUT2	M-trgovina (Rumunjska)	PE, HM, SI, TR, C-19	Rad koristi iznimno izmijenjen UTAUT2 koji je od standardnih faktora ostavio samo PE, HM i SI. Izostavljen je i faktor USE, a dodani su Povjerenje – TR (<i>Trust</i>) i

				BI	utjecaj bolesti Covid-19 na namjeru ponašanja kod korištenja M-trgovine.
68.	Xian (2021)	UTAUT2	Primjena umjetne inteligencije (Kina)	PE, EE, SI, FC, HM, PV, HT, PI BI	Rad koristi UTAUT2 model iz kojeg je izostavljena zavisna varijabla USE, a umjesto standardnih moderatora (spol, dob i iskustvo) koristi se Osobna inovativnost – PI (<i>Personal innovativeness</i>) koja je zamišljena da moderira utjecaj preostalih sedam varijabli na namjeru ponašanja.

Izvor: izrada autora

***** PE – Performance Expectancy, EE – Effort Expectancy, SI – Social influence, FC – Facilitating conditions, PV – Price value, HM – Hedonic Motivation, HT – Habit, BI – Behavioral intention**

Podebljanim slovima označeni su konstrukti koji su dodani i nisu dio standardnog UTAUT2 modela.

**** ICAAM - Integrated chatbot acceptance-avoidance model**

**** UTAUT4-AV – Unified Theory of Acceptance and Use for Automated Vehicles**

**** BDI – Belief Desire Intentions**

**** ISS - Information systems success**

U tablici 9. za pregled literature nisu korišteni samo radovi koji istražuju prihvaćanje i korištenje tehnologije kroz UTAUT2 model već i kroz druge modele. Shodno tome, nije za očekivati kako će svi radovi navedeni u tablici 9. sadržavati barem osnovne konstrukte modela UTAUT2. Od ukupno 68 radova koliko je analizirano u tablici 9., njih 51 (75 %) koristi UTAUT2 model, bilo da istražuju na UTAUT2 modelu samostalno, bilo da ga koriste u kombinaciji s drugim modelom ili je model u manjoj ili većoj mjeri modificiran, no autori ga ipak nazivaju UTAUT2 modelom. Neki od preostalih korištenih modela u radu su ICAAM, UTAUT4AV, UTAUT, TAM, BRT, AIDUA, GAAIS, UTAUT3, TPB i Triandisov model, a neki od modela koji su korišteni u kombinaciji s UTAUT2 modelom u pregledu literature su SDT, TTF, DOI, BDI i ISS.

Kako se navodi u tablici 10., od ukupno 68 radova koji su analizirani u pregledu literature, njih 93 % koristi konstrukt PE, 91 % koristi konstrukt EE, 88 % koristi SI, 84 % koristi konstrukt FC, 72 % koristi konstrukt HM, 62 % koristi konstrukt HT, a samo 49 % pregledanih radova koristi konstrukt PV.

Osim navedenoga važno je napomenuti kako 87 % pregledanih radova koristi faktor namjere korištenja (BI - *Behavioral intention / Intention to use*), a samo 44 % radova koristi faktor stvarnog korištenja tehnologije (*Use behavior / USE / Actual usage*).

Tablica 10. Korišteni standardni konstrukti iz UTAUT2 modela

Oznaka	Konstrukt	Izvorni model	Broj radova	Postotak
PE	<i>Performance Expectancy</i> Očekivani učinak	UTAUT	63	93 %
EE	<i>Effort Expectancy</i> Očekivani napor	UTAUT	62	91 %
SI	<i>Social influence</i> Društveni utjecaj	UTAUT	60	88 %
FC	<i>Facilitating conditions</i> Olakšavajući uvjeti	UTAUT	57	84 %
HM	<i>Hedonic motivation</i> Hedonistička motivacija	UTAUT2	49	72 %
PV	<i>Price value</i>	UTAUT2	33	49 %

	Cijena			
HT	<i>Habit</i> Navika	UTAUT2	42	62 %

Izvor: izrada autora

Osim sedam standardnih konstrukta navedenih u tablici 10., u pregledu literature modeli su često prošireni i nadopunjeni s novim konstruktima, ukupno više od 80 konstrukta koji su spojeni u određene kategorije. Primjerice, konstrukti povjerenje (TR – *Trust*), potrošačko povjerenje (CT – *Consumer trust*), percipirano povjerenje (PT – *Perceived trust*) spojeni su pod jedan krovni konstrukt vezan za povjerenje. Jednako tako učinjeno je s konstruktima povezanim sa sigurnošću i rizicima te konstruktima vezanim za privatnost, konstruktima vezanim za osobnu inovativnost, percipiranu kvalitetu, održivost i vrijednost proizvoda.

Konstrukti koji se spominju u manjoj mjeri su percipirana prijetnja, percipirana ozbiljnost, osjetljivost, praktična udobnost, zdravlje, ugodnost, percipirana lakoća korištenja, razlozi za prihvaćanje ili odbijanje, barijere u korištenju, optimizam u vezi umjetne inteligencije, percipirana autonomnost, emocije, spremnost na korištenje, spremnost na plaćanje, moralna odgovornost, percipirani ljudski dodir, dizajn, etika, kompatibilnost, radoznalost, percipirane prednosti proizvoda, vladina pomoć i društvena prisutnost.

Tablica 11. Najčešće korišteni konstrukti kao proširenja modela za istraživanje prihvaćanja i korištenja tehnologije

Kod	Konstrukt	Broj radova	Postotak
S-R	<i>Security and risks</i> Sigurnost i rizici	21	31 %
TR	<i>Trust</i> Povjerenje	15	22 %
PI	<i>Personal innovativeness</i> Osobna inovativnost	12	18 %
PP	<i>Privacy</i> Privatnost	10	15 %
ATT	<i>Attitude</i> Stav	8	12 %
PQ	<i>Perceived quality</i> Percipirana kvaliteta	8	12 %

PU	<i>Perceived usefulness</i> Percipirana korisnost	4	6 %
PV	<i>Product value</i> Vrijednost proizvoda	4	6 %
S	<i>Sustainability</i> Održivost	3	4 %
ANTH	<i>Anthropomorphism</i> Antropomorfizam	3	4 %
KN	<i>Knowledge or skills</i> Znanje i vještine	3	4 %
PI	<i>Perceived interactivity</i> Percipirana interaktivnost	3	4 %
ANX	<i>Anxiety</i> Anksioznost	3	4 %

Izvor: izrada autora

4.1. Modifikacije i proširenja UTAUT2 modela

4.1.1. Stvarno korištenje (USE)

S obzirom na to da izvorni rad u kojem Venkatesh i sur. (2012) predstavljaju UTAUT2 model ima vrlo specifične čestice za faktor korisničkog ponašanja kod stvarnog korištenja tehnologije (USE), a koji nije primjenjiv na mnogobrojnim istraživanjima, autori su u prilagođenim verzijama navedenog modela koristili različite čestice kako bi skrojili konstrukt USE te mjerili ponašanje korisnika kod korištenja tehnologije.

Venkatesh i sur. (2012) u svom su radu koristili UTAUT2 model u kontekstu prihvatanja i korištenja mobilnog interneta, a konstrukt koji je mjerio stvarno korištenje tehnologije imao je samo jedno pitanje s mogućim višestrukim odgovorima:

Molimo vas odaberite učestalost korištenja za svaki od:

- a) SMS
- b) MMS
- c) skidanje multimedije
- d) Java igre
- e) pretragu internetskih stranica

f) mobilna e-pošta

*Učestalost korištenja izražava se na Likertovoj ljestvici od 1 do 7 pri čemu je 1 „nikada“, a 7 „puno puta u danu“.

Problem nastaje kod potrebe istraživača da istražuje korištenje određene tehnologije za koju nije potpuno izvedivo nabrojati određene potkategorije i mjeriti korištenje tehnologije na navedeni način. U tablici 12. prikazane su varijacije čestica za mjerenje konstrukta USE, odnosno kako su pregledana istraživanja prilagođavala navedeni konstrukt za kontekst svog istraživanja.

Tablica 12. Varijacije čestica za konstrukt stvarnog korištenja

Autori	Kontekst	Čestice
Ameri i sur. (2020)	Mobilna aplikacija za obrazovanje u farmaciji (LabSafety)	USE1: Redovito koristim aplikaciju LabSafety. USE2: Korištenje aplikacije LabSafety je ugodno iskustvo. USE3: Trenutno koristim aplikaciju LabSafety. USE4: Provodim puno vremena na aplikaciji LabSafety.
Baptista & Oliveira (2015)	Usluge mobilnog bankarstva	Kolika je vaša stvarna učestalost korištenja usluga mobilnog bankarstva? (1) ne koristim; (2) jednom godišnje; (3) jednom u šest mjeseci; (4) jednom u tri mjeseca; (5) jednom mjesečno; (6) jednom tjedno; (7) jednom u 4-5 dana; (8) jednom u 2-3 dana; (9) gotovo svakodnevno; (10) svaki dan; (11) nekoliko puta dnevno
Budhathoki i sur. (2024)	ChatGPT	UB1: Koristim besplatnu verziju ChatGPT-a (UB1). UB2: Koristim ChatGPT kao pomoćnika u pisanju pogonjenog umjetnom inteligencijom. UB3: Koristim ChatGPT za izradu ocijenjenih radova.
Cabrera-Sanchez i sur. (2021)	Primjena umjetne inteligencije	UB1: Karte - Koliko trenutno koristite karte i rute? UB2: Preporuke – Koliko trenutno koristite preporuke? UB3: Glas – Koliko trenutno koristite glasovne naredbe?
Chatterjee & Bhattacharjee (2020)	Primjena umjetne inteligencije u visokom obrazovanju	Primjena umjetne inteligencije u visokom obrazovanju dobra je za društvo.

		Primjena umjetne inteligencije u visokom obrazovanju učinit će obrazovanje interaktivnijim. Primjena umjetne inteligencije u visokom obrazovanju učinit će ga troškovno isplativim. Primjena umjetne inteligencije u visokom obrazovanju učinit će aktivnosti podučavanja i učenja zanimljivijim.
Chopdar i sur. (2018)	Aplikacije za m-trgovinu	Navedite u kojoj mjeri se slažete s navedenim tvrdnjama: (1 – uopće se ne slažem, 7 – u potpunosti se slažem) UB1: U posljednjih 6 mjeseci kupovao/la sam putem mobilnih <i>shopping</i> aplikacija s namjerom kupnje online proizvoda. UB2: U posljednjih 6 mjeseci kupovao/la sam putem mobilnih <i>shopping</i> aplikacija s namjerom kupnje proizvoda od različitih online prodavača. UB3: U posljednjih 6 mjeseci kupovao/la sam putem mobilnih <i>shopping</i> aplikacija s namjerom da napravim osobnu kupnju. UB4: Koristio/la sam nekoliko različitih mobilnih <i>shopping</i> aplikacija u posljednjih 6 mjeseci.
Gansser & Reich (2021)	Primjena umjetne inteligencije (za mobilnost / kućanstvo / zdravlje)	Potkonstrukti su mjereni po uzoru na Venkatesh i sur. (2012) s mjerenjem učestalosti korištenja na temelju 6 čestica. Čestice za mobilnost: Pomoćnik u vožnji; Autonomna vožnja; Informacijski i navigacijski sustav; Detekcija stanja vozača; Kontrola prometa; Prediktivno održavanje Čestice za kućanstvo: Robotski usisavač; <i>Shopping</i> / Kuhinjski pomagač; Inteligentni sustavi za pametnu kuću; Pametni sat; Glasovni asistenti (npr. Alexa); Inteligentni <i>chatbotovi</i> za kupnju ili korisničku službu Čestice za zdravlje: Inteligentni osobni asistenti za osobu s posebnim potrebama; Roboti za srčb (npr. za kretanje); Pametni uređaji za praćenje zdravlja; Pametni uređaji u

		bolnicama; Sustavi za pomoć u svakodnevnim životnim aktivnostima; Pametni sustavi za savjetovanje u dijagnosticiranju i preporuci terapije
Kilani i sur. (2023)	Korištenje digitalnih novčanika / e-novčanika	UB1: Obično koristim svoj e-novčanik više puta tijekom određenog mjeseca. UB2: Koristim svoj e-novčanik za veliki dio svojih financijskih transakcija u određenom mjesecu. UB3: Iskorištavam većinu dostupnih upotreba svog e-novčanika (plaćanja računa, podizanje gotovine, kupnje na prodajnom mjestu i <i>online</i> kupovinu).
Lavidas i sur. (2024)	Korištenje aplikacija umjetne inteligencije (ChatGPT)	UB1: Namjeravam nastaviti koristiti aplikacije umjetne inteligencije kao što je ChatGPT u svojim studijima. UB2: Uvijek ću nastojati koristiti aplikacije umjetne inteligencije kao što je ChatGPT u svojim studijima. UB3: Planiram često koristiti aplikacije umjetne inteligencije kao što je ChatGPT u svojim studijima.
Limayem & Hirt (2003)	Interaktivna internetska komunikacijska tehnologija (WebBoard)	AB1: Koliko ste puta tjedno tijekom prošlog mjeseca pristupili WebBoardu? AB2: Koliko ste poruka tjedno objavili na WebBoardu tijekom prošlog mjeseca? <i>(Na pitanje se odgovara na ljestvici od 1 do 7 pri čemu je 1 – Nisam uopće, a 7 – Nekoliko puta u danu)</i>
Macedo (2017)	Informacijsko-komunikacijska tehnologija	Vrijeme korištenja interneta: (1) kraće od 1 godine (2) 1 do 2 godine (3) duže od 2 godine Učestalost korištenja interneta: (1) svaki dan ili skoro svaki dan (2) barem jednom tjedno (3) barem jednom mjesečno (4) rjeđe od jednom mjesečno Korištene internet aktivnosti: (1) Slanje i primanje mailova; (2) Traženje informacija; (3) Rezervacije (npr. za hotel ili restoran); (4) Traženje zdravstvenih savjeta; (5) Vijesti i opće informacije; (6)

		Internetska kupovina; (7) Financijske aktivnosti; (8) Facebook; (9) Chatovi
Maruping i sur. (2017)	Informacijska tehnologija	Maruping i sur. (2017) navode kako su za varijablu USE koristili čestice po uzoru na Venkatesh i sur. (2008) – učestalost korištenja, dužina korištenja i intenzitet korištenja.
Venkatesh i sur. (2008)	Informacijski sustavi	Trajanje: U prosjeku, koliko sati tjedno koristite sustav? Učestalost: Koliko često koristite sustav? Intenzitet: Kako procjenjujete opseg vaše trenutne uporabe sustava?
Mohd Rahim i sur. (2022)	Korištenje <i>chatbota</i> kod institucija za visoko obrazovanje	USE1: Spreman sam ponovno koristiti <i>chatbot</i> u budućnosti. USE2: Slažem se da će korištenje <i>chatbota</i> poboljšati moje iskustvo u rješavanju mojih akademskih pitanja. USE3: Slažem se da je <i>chatbot</i> koristan za besprijekorno rukovanje čestim akademskim pitanjima i odgovorima (FAQ). USE4: Upotreba <i>chatbota</i> moći će podržati moja nenadzirana i hitna pitanja vezana za akademski život.
Nikolopoulou i sur. (2020)	Korištenje mobilnih telefona u svrhu studija	USE1: Redovito koristim mobilni telefon u svrhu studija. USE2: Korištenje mobilnog telefona ugodno je iskustvo. USE3: Trenutno koristim mobilni telefon kao podržavajući alat za studiranje. USE4: Tijekom studija puno vremena provodim koristeći mobilni telefon.
Sobaih i sur. (2024)	Korištenje ChatGPT-a u visokom obrazovanju	AU1: Namjeravam koristiti znanja i vještine koje sam stekao/la koristeći ChatGPT u svojim obrazovnim aktivnostima. AU2: Znanje i vještine koje sam stekao/la koristeći ChatGPT bit će mi korisni u nastavi. AU3: Korištenje ChatGPT-a mi je pomoglo poboljšati akademske rezultate.

Stzelecki (2023)	Korištenje ChatGPT-a u visokom obrazovanju	Molimo Vas odaberite učestalost svog korištenja ChatGPT-a: (1) Nikada; (2) Jednom mjesečno; (3) Nekoliko puta mjesečno; (4) Jednom tjedno; (5) Nekoliko puta tjedno; (6) Jednom dnevno; (7) Nekoliko puta dnevno
Tseng i sur. (2022)	Sustavi za online tečajeve (MOOC)	Tjedna dužina korištenja sustava za <i>online</i> tečajeve (uključujući i pripremu za nastavu) Tjedna učestalost korištenja sustava za <i>online</i> tečajeve (uključujući i pripremu za nastavu)
Vimalkumar i sur. (2021)	Virtualni glasovni asistenti	Učestalost: Kada ste zadnji put na mobilnom uređaju koristili neku od mogućnosti virtualnih glasovnih asistenata? (1) Nikada (2) Prije duže od 90 dana (3) U zadnjih 30 dana (4) U zadnjih 7 dana (5) Danas Dodatak 1: Alarm Dodatak 2: Podsjetnik Dodatak 3: Vrijeme Dodatak 4: Vijesti

Izvor: izrada autora

U preliminarnom istraživanju, Biloš i Budimir (2024) koristili su čestice za konstrukt USE prilagođene prema Nikolopoulou i sur. (2020) što se nije pokazalo uspješnim zbog visoke multikolinearnosti između BI i USE. Nakon analize postojećih varijacija u mjerenju konstrukta stvarnog korištenja (USE), a uzevši u obzir specifičnost generativne umjetne inteligencije kao tehnologije, odabran je pristup koji navode Venkatesh i sur. (2008). Navedeni pristup mjerenja korištenja podrazumijeva ispitivanje trajanja korištenja (U prosjeku, koliko sati tjedno koristite sustav?), učestalosti korištenja (Koliko često koristite sustav?) i intenziteta korištenja (Kako procjenjujete opseg Vaše trenutne upotrebe sustava?). Mjerne ljestvice za mjerenje trajanja, učestalosti i intenziteta su prilagodljive kontekstu istraživanja te ih autor treba skrojiti shodno tome. Naime, ukoliko bi se mjerilo korištenje pametnih telefona i korištenje tehnologije za pametnu kuću, nije moguće uspoređivati obujam korištenja takvih tehnologija. Upravo iz tog razloga navedene mjerne ljestvice moraju biti adaptivne kontekstu istraživanja.

4.1.2. Osobna inovativnost – (PI *Personal innovativeness*)

Koncept osobne inovativnosti (PI – *Personal innovativeness*) skrojili su i unutar istraživanja predstavili Agarwal i Prasad (1998). Navedeni autori dodaju kako se koncept osobne inovativnosti odnosi na stupanj u kojem je pojedinac spreman isprobati nove komunikacijske tehnologije te zaključuju kako određeni pojedinci usvajaju nove tehnologije ranije od drugih. Osim toga, kako navode Gansser i Reich (2021), pretpostavlja se kako ljudi s izraženijom osobnom inovativnošću posjeduju ili koriste više tehnoloških uređaja i inovacija jer su kroz veću interakciju s tehnologijom stekli pozitivno iskustvo s tehnologijom te im to daje osjećaj kako će i novi proizvod dobro funkcionirati pa su zbog toga spremniji prihvaćati nove tehnološke inovacije. Bruschi i Rappell (2020) navode kako je znatiželja i sklonost učenju o novim tehnologijama obilježje inovativnih pojedinaca, stoga je onda osobna inovativnost pozitivno povezana s namjerom korištenja nove tehnologije. Huntley i Chacko (2012) zaključuju kako su pripadnici Generacije Y, poznati i kao milenijalci, visokoobrazovana i tehnološki potkovana generacija koja vjeruje kako je tehnologija integrirana u njihov život (Xian, 2021; Oliver, 2007). Strzelecki (2023) navodi kako osobna inovativnost ima pozitivan učinak na namjeru Generacije Z da koriste ChatGPT, alat generativne umjetne inteligencije, a isti zaključak navode i Biloš i Budimir (2024). Može se zaključiti kako je kod mlađih generacija osobna inovativnost značajan prediktor u namjeri ponašanja. S druge strane, podatci pokazuju kako starija populacija nije osobito sklona inovacijama ukoliko ne moraju (Lee i sur., 2010; Tams & Dulipovici, 2024).

Važnost koncepta osobne inovativnosti leži u činjenici da inovativni pojedinci imaju ključnu ulogu u širenju inovacija te su zbog toga cijenjeni na tržištu (Senali i sur., 2023). Garcia de Blanes Sebastian i sur. (2022) te Gansser i Reich (2021) sugeriraju da osobna inovativnost utječe na namjeru ponašanja, te da je iznimno važna kod istraživanja novih tehnologija kao što je umjetna inteligencija. Xian (2021) navodi kako se osobna inovativnost može integrirati u UTAUT2 model i kao moderirajući faktor, kako su ga i u izvornom radu koristili Agarwal i Prasad (1998) u svom konceptualnom modelu. Morgan i Morgan (2006) u svom istraživanju ističu kako muškarci općenito iskazuju pozitivnije stavove prema tehnologiji nego li to iskazuju žene, a što u konačnici rezultira i većoj stopom usvajanja novih tehnologija.

Uobičajene čestice konstrukta osobne inovativnosti su (Gansser i Reich, 2021; Biloš i Budimir, 2024; Foroughi i sur., 2023):

- Kada čujem za novu tehnologiju potražim način da ju odmah isprobam.

- Među mojim prijateljima, ja sam uglavnom prvi/a koji proba nove tehnologije.
- Volim eksperimentirati s novim tehnologijama.

Uz navedene tri čestice, konstrukt se ponekad dodaje i čestica „Općenito, oklijevam s isprobavanjem nove tehnologije“; ali se kod navedene čestice radi obrnuta mjerna ljestvica (Baudier i sur., 2020; Garcia de Blanes Sebastian i sur., 2022). Za potrebe ovoga rada konstrukt osobne inovativnosti koristit će tri čestice. Za utjecaj osobne inovativnosti na namjeru ponašanja skrojena je hipoteza koja glasi:

- *Utjecaj osobne inovativnosti na namjeru ponašanja bit će moderiran dobi i spolom, tako da će efekt biti jači za mlađu populaciju, a posebno mlađe muškarce.*

4.1.3. Sigurnost i rizici

Kako je ranije spomenuto, konstrukti vezani za sigurnost i percipirane rizike korištenja tehnologije spominju se u gotovo trećini pregledanih radova, a među njima ipak je najčešći konstrukt percipiranih rizika. Osim toga, konstrukti koji su vezani za percipirane rizike međusobno su si sličniji nego li konstrukti vezani za sigurnost. Osim međusobne sličnosti, navedeni konstrukti imaju čestice koje se naslanjaju na srodne konstrukte koji su također navedeni u tablici 10., prvenstveno misleći na povjerenje, privatnost, kvalitetu proizvoda pa čak i anksioznost. Čestice konstrukta sigurnosti koje navode Gansser i Reich (2021) spoj su konstrukta anksioznosti, privatnosti i sigurnosti:

- Zabrinut/a sam za svoje osobne podatke prilikom korištenja proizvoda u MO/HH/HA koji sadrže umjetnu inteligenciju (AI).
- Zabrinut/a sam za sigurnost podataka proizvoda u MO/HH/HA koji sadrže umjetnu inteligenciju (AI).
- Zabrinut sam za privatnost kod proizvoda u MO/HH/HA koji sadrže umjetnu inteligenciju (AI).
- Zabrinut sam za sigurnost kod proizvoda u MO/HH/HA koji sadrže umjetnu inteligenciju (AI).

Postoje i drugi pristupi u kojima konstrukt sigurnosti kombinira s konstrukom povjerenja, ali više naginje na konstrukt povjerenja (TR – *Trust*) nego li na sigurnosni konstrukt (Korkmaz i sur., 2022):

- Mislim da je APTS pouzdan.
- Mislim da je APTS siguran.

- Općenito, mogu vjerovati APTS-u.
- Mislim da je APTS sigurniji od tradicionalnog javnog prijevoza.
- Mislim da bi APTS pomogao smanjiti prometne nesreće.

Unutar istoga rada, u kojem se istražuje prihvaćanje i korištenje autonomnih sustava javnog prijevoza, autori navode i konstrukt pod nazivom percipirani rizik, a čestice navedenog konstrukta su sljedeće: (Korkmaz i sur.,2022):

- Korištenje ATPS sustava dovelo bi do financijskog gubitka za mene.
- APTS sustav možda neće dobro funkcionirati i stvorit će probleme.
- Korištenje APTS sustava bilo bi rizično.

Primjeri čestica za konstrukt percipiranog rizika koje navode Lai i sur. (2024) kontekstom su primjenjiviji na istraživanjima vezanima za alate generativne umjetne inteligencije, premda se i ovdje može uočiti kako se ovaj konstrukt kontekstom naslanja na konstrukte vezane za privatnost:

- Dobio/la bih kaznu za plagijat ako koristim ChatGPT za obavljanje zadataka.
- Ako bih koristio/la ChatGPT za obavljanje zadataka, vjerojatno bih bio/la uhvaćen/a.
- Mogao/la bih biti razočaran/a ako obavim zadatke uz pomoć ChatGPT-a.
- Mislim da je korištenje ChatGPT-a za obavljanje zadataka ugrožava moju privatnost.

Čestice koje navode Chatterjee i Bhattacharjee (2020), a čije se istraživanje također dotiče primjene umjetne inteligencije, također se naslanja na druge srodne konstrukte kao što su povjerenje i točnost ili kvaliteta proizvoda, te navode sljedeće čestice:

- Vjerujem da obrazovani sadržaj generiran od strane umjetne inteligencije nije uvijek točan.
- Primjena umjetne inteligencije u svrhu upisa je zbunjujuća.
- Radije ne bih koristio/la umjetnu inteligenciju u administrativne svrhe.
- Korištenje umjetne inteligencije za odgovaranje na studentske upite je rizično.

Percepcija sigurnosti ili percipirana sigurnost definira se kao stupanj vjerovanja i povjerenja, a sigurnosni propusti smatraju se značajnom preprekom koja sprječava potrošače u namjeri korištenja određene tehnologije (Merhi i sur., 2019). Iz navedene se teorije može vidjeti kako

je konstrukt percipirane sigurnosti vrlo srodan konstruktima povjerenja, rizika, ali i konstrukt koji se veže za zabrinutost po pitanju privatnosti. Navedeni konstrukti iznimno su bitni za istraživanja u kontekstu bankarskih sustava, medicinskih sustava, sustava autonomnih vozila i drugih tehnologija gdje postoji realna bojazan od zloupotrebe sustava (Merhi i sur., 2019). Percipirani rizici (PR – *Perceived risks*) definiraju se kao percepcija pojedinaca o nesigurnosti i ozbiljnosti posljedica povezanih s određenim ponašanjem (Lai i sur., 2024). Hanafizadeh i sur. (2014) navode kako visok rizik od korištenja tehnologije negativno utječe na namjeru pojedinca da ju koristi. Osim kod navedenih sustava kod kojih su konstrukti vezani za sigurnost i rizike iznimno važni, ovaj konstrukt može biti važan i za alate generativne umjetne inteligencije jer se pokazalo kako ChatGPT lažno generira pregled akademske literature, daje izmišljene citate i koristi nepostojeće reference, a što može biti rizik za studente koji koriste ChatGPT za svoje studentske zadatke (Lai i sur., 2024; Rahman i sur., 2023). S druge strane, navedena problematika može se testirati i s konstruktima vezanima za percipiranu kvalitetu tehnologije ili percipiranu korisnost. Sun i sur. (2023) također navode novitet i koriste znanje kao moderator. Njihovo je istraživanje pokazalo kako su kod visoke razine znanja o tehnološkom proizvodu ili usluzi inovativnost i navike značajnije povezane s korištenjem. S druge strane, potrošači s niskom razinom znanja bit će pod većim utjecajem društvenog utjecaja.

Uzevši u obzir prikazane čestice koje se koriste za konstrukt sigurnosti i rizika te kontekste u kojima se najčešće koriste spomenute ekstenzije na UTAUT2 model, u ovom radu ne postoji potreba za dodavanjem navedenog konstrukta u istraživanje o prihvatanju i korištenje generativne umjetne inteligencije, premda je to najčešće integrirana ekstenzija koja se spominje u pregledu literature.

4.1.4. Povjerenje (TR – *Trust*)

Povjerenje se definira kao osjećaj sigurnosti i spremnosti koju osoba stavlja u sustav, uslugu ili proizvod koji redovito ispunjava njihova očekivanja (Kim i sur., 2017). Da bi se tehnologija smatrala pouzdanom i vjerodostojnom, od nje se očekuje da nadvlada skepticizam (Salifu i sur., 2024). Povjerenje se shvaća kao multidimenzionalni koncept koji odražava percepcije kompetentnosti, integriteta i dobronamjernosti druge strane (Mayer i sur., 1995). Kada je tehnologija nova i nepoznata korisnicima, oni se osjećaju nesigurno, a povjerenje je jedan od najvažnijih elemenata za prevladavanje nesigurnosti (Garcia de Blanes Sebastian i sur., 2022). Povjerenje je ključni element u pomaganju pojedincima da prevladaju zabrinutost i strah od rizika vezanih za gubitak privatnosti povezanih s korištenjem tehnologije, posebno one

internetski povezane (Dinev i sur., 2016; Vimalkumar i sur., 2021). U istraživanju u kojem istražuju prihvaćanja i korištenje virtualnih glasovnih asistenata, Vimalkumar i sur. (2021) zaključuju kako konstrukti Zabrinutost za privatnost i Percipirani rizici vezani za privatnost nisu potvrđeni kao prediktori namjere korištenja tehnologije, dok se konstrukt povjerenja potvrdio statistički značajnim. Shi i sur. (2022) smatraju kako će IoT (engl. *Internet of Things*; hrv. Internet stvari) biti široko prihvaćen ako ljudi vjeruju u njegovu sposobnost da ispuni svoja obećanja i ukoliko vjeruju u tvrtku koja pruža tu uslugu ili proizvodi tu tehnologiju, a isto se može reći i za generativnu umjetnu inteligenciju. Kako navode Lai i sur. (2024), kod korištenja alata generativne umjetne inteligencije kao što je ChatGPT, povjerenje može biti snažan motivator namjere ponašanja za korištenje velikih jezičnih modela i sličnih tehnologija, ali s obzirom na to da generirani rezultati nekada budu netočni i pristrani, u izostanku alternative koja će procijeniti pouzdanost generiranih rezultata, nedostatak povjerenja može negativno utjecati na namjeru korištenja takvih tehnologija. Istraživanje koje su proveli Mohd Rahim i sur. (2022) ukazuje na to kako je percipirano povjerenje prediktor namjere ponašanja kod korištenja *chatbotova* pogonjenih umjetnom inteligencijom. Utjecaj povjerenja na namjeru ponašanja potvrdilo je i istraživanje koje su proveli Mehedi i sur. (2023) u kontekstu prihvaćanja i korištenja ChatGPT-a kao alata generativne umjetne inteligencije.

Mlađi ljudi iskusniji su s tehnologijom pa su općenito pozitivnijeg stava prema tehnologiji, pa tako i prema umjetnoj inteligenciji i njezinom utjecaju na čovječanstvo, dok su stariji ispitanici skloniji vidjeti tako naprednu tehnologiju kao prijetnju ljudima (Broady i sur., 2010; Lee i sur., 2017; Elias i sur., 2012; Schepman i Rodway, 2023). Mays i sur. (2020) navode kako su mlađi ispitanici, posebno mlađi muškarci otvoreniji po pitanju stava o korištenju umjetne inteligencije. S druge strane, iskustvo je povezano s osjećajem anksioznosti kod korištenja umjetne inteligencije, a što u konačnici znači da će s većim iskustvom korisnik postati opušteniji u korištenju tehnologije (Nomura i sur., 2006; Mays i sur., 2020). Upoznatost s tehnologijom smanjuje neizvjesnost i otklanja zabrinutost kod korištenja tehnologije. Veća razina upoznatosti vodi do većeg znanja temeljem prethodnih iskustava, a što rezultira većim povjerenjem u tehnologiju (Gefen, 2000; Choung i sur., 2023). Helberger i sur. (2020) upućuju na pozitivnu korelaciju između korisničkog iskustva i povjerenja u smislu stava korisnika prema tehnologijama umjetne inteligencije.

Garcia de Blanes Sebastian i sur. (2022) provodili su istraživanje o prihvaćanju i korištenju virtualnih glasovnih asistenata kao što je primjerice Alexa, a čestice koje navode pod konstruktom povjerenja su sljedeći:

- TR1: Uređaji s glasovnim asistentom su pouzdani.
- TR2: Vjerujem uređajima s glasovnim asistentom zbog njihove sposobnosti da izvrše svoje funkcije.
- TR3: Uređaji s glasovnim asistentom su sposobni izvršiti zadane zadatke.
- TR4: Uređaji s glasovnim asistentom još uvijek imaju moje povjerenje.

Korkmaz i sur. (2022), koji su proveli istraživanje o prihvaćanju i korištenju autonomnih sustava za javni prijevoz (APTS), koriste kontrukst pod nazivom Povjerenje i sigurnost, a njegove čestice su sljedeće:

- TS1: Mislim da je APTS pouzdan.
- TS2: Mislim da je APTS siguran.
- TS3: Općenito, mogu vjerovati APTS-u.
- TS4: Mislim da je APTS sigurniji od tradicionalnog javnog prijevoza.
- TS5: Mislim da bi APTS pomogao smanjiti prometne nesreće.

Lai i sur. (2024), koji su provodili istraživanje o prihvaćanju i korištenju ChatGPT-a, za mjerenje konstrukta povjerenja koriste iduće čestice:

- TR1: Vjerujem da je ChatGPT pouzdan za izvršavanje mojih zadataka.
- TR2: Vjerujem u kvalitetu odgovora ChatGPT-a u izvršavanju mojih zadataka.
- TR3: Uvjeren sam da je pružatelj tehnologije ChatGPT pošten.
- TR4: Čak i ako nije nadzirano, vjerovao bih da bi zadatci mogli biti ispravno izvršeni s funkcijama ChatGPT-a.

Mohd Rahim i sur. (2022) proveli su istraživanje o prihvaćanju i korištenju *chatbotova* pogonjenih umjetnom inteligencijom, a čestice koje su koristili za konstrukt percipiranog povjerenja su sljedeće:

- PT1: Koristit ću *chatbot* ako osjetim da je sadržaj pouzdan.
- PT2: Koristit ću *chatbot* ako osjetim da chatbot pruža pouzdane informacije.
- PT3: Koristit ću *chatbot* ako osjetim da chatbot ispunjava moja očekivanja.
- PT4: Koristit ću *chatbot* ako osjetim da je chatbot siguran.

Čestice za konstrukt povjerenja koje navode Vimalkumar i sur. (2021) u istraživanju o prihvaćanju i korištenju virtualnih glasovnih asistenata su sljedeći:

- TR1: Virtualni glasovni asistenti su sigurna okruženja za razmjenu informacija s drugima.
- TR2: Virtualni glasovni asistenti su pouzdana okruženja za obavljanje poslovnih transakcija.
- TR3: Virtualni glasovni asistenti kompetentno rukuju osobnim informacijama koje korisnici dostavljaju.
- TR4: Mislim da su virtualni glasovni asistenti pouzdani.
- TR5: Osjećam se sigurno da me pravne i tehnološke strukture adekvatno štite od problema s virtualnim glasovnim asistentima.

Uzevši u obzir sve navedeno, zaključak je autora kako postoji potreba za proširenjem UTAUT2 modela s konstruktom povjerenja kada je riječ o istraživanju u kontekstu prihvatanja i korištenja generativne umjetne inteligencije. Ukoliko se kao najbolja verzija čestica za konstrukt povjerenja uzme ona koju navode Mohd Rahim i sur. (2022), onda je dovoljno imati samo taj konstrukt jer on objedinjuje i povjerenje i rizike, ali i percipiranu razinu kvalitete, korisnosti i točnosti, što su također konstrukti s liste najzastupljenijih ekstenzija UTAUT2 modela. Dakle, u slučaju ovoga rada dodan je konstrukt povjerenje koji objedinjuje nekoliko različitih konstrukta, a hipoteza vezana za ovaj konstrukt glasi:

- *Dob, spol i iskustvo će moderirati učinak povjerenja na namjeru ponašanja, tako da će učinak biti jači među mlađom populacijom, posebno mlađim muškarcima u kasnijim fazama korištenja tehnologije.*

4.2. Moderator

Od ukupno 68 radova koliko je uzeto u razmatranje u okviru pregleda literature unutar ovoga rada, tek 26 radova (38 %) u okviru svog istraživanja koristi moderatore. Dakle, niti polovica radova nije koristila jedan od najvažnijih elemenata UTAUT2 modela. Venkatesh i sur. (2012) navode tri standardna moderatora unutar UTAUT2 modela – dob, spol i iskustvo. U odnosu na UTAUT, izbačen je moderator dobrovoljnosti jer je UTAUT2 koncipiran da se na njemu provode istraživanja na krajnjim potrošačima tehnologije.

Ameri i sur. (2019) navode kako su u istraživanju izostavili moderator iskustva jer ispitanici nisu mogli ranije iskusiti tehnologiju koja je u središtu njihova istraživanja. Stoga su u istraživanju koristili dva moderatora – godine i spol. Godine nisu imale statistički značajan učinak niti na jedan konstrukt, dok se kod spola pokazalo kako je navika na korisničko

ponašanje snažnije utjecala kod muškaraca nego li kod žena. Istraživanje koje su proveli Maican i sur. (2023) zanimljivo je i iz razloga što kao moderatore koriste kreativnost i razinu znanja engleskog jezika. Maican i sur. (2023) navode kako očekivani naponi imaju snažan pozitivan učinak na namjeru ponašanja kod osoba s niskom kreativnošću, ali je zanimljivo i kako kreativnost smanjuje pozitivan odnos između očekivane izvedbe i namjere ponašanja te navike i namjere ponašanja. Poznavanje engleskog jezika posebno je zanimljiv i bitan moderator za istraživanja koja se provode na ispitanicima kojima engleski nije materinski jezik. Maican i sur. (2023) navode kako poznavanje engleskog jezika jača pozitivan utjecaj društvenog utjecaja na namjeru ponašanja. Navedeno istraživanje pokazalo je kako je utjecaj društvenog utjecaja na namjeru ponašanja kod sudionika sa srednjim ili naprednim znanjem engleskog jezika jači nego kod onih koji slabije znaju engleski jezik. Osim toga, istraživanje je pokazalo kako moderatori predstavljaju važno povećanje objašnjenja varijance zavisne varijable namjere ponašanja. Maruping i sur. (2016) u svom istraživanju potvrđuju teze koje navode Venkatesh i sur. (2003), a koje upućuju na to kako su očekivani naponi snažan prediktor namjere ponašanja, posebice kod starijih žena, dok je očekivana izvedba snažan prediktor namjere ponašanja za mlade muškarce. Osim navedenih moderatora, kako je ranije navedeno kod konstrukta osobne inovativnosti, PI se može koristiti i kao konstrukt, ali i moderator (Shaw & Sergueeva, 2019; Xian, 2021; Foroughi i sur., 2023). S druge strane, Shi i sur. (2022) navode kako i konstrukt olakšavajućih uvjeta može biti korišten kao moderator te je njihovo istraživanje i pokazalo kako veća očekivanja izvedbe rezultiraju većom spremnošću za prihvatanjem IoT tehnologije ako je korisnik opremljen s boljim uvjetima koji mu olakšavaju korištenje, kao što je primjerice dobra internetska infrastruktura, instrumentalna podrška i slično. Točnost informacija još je jedan u nizu moderatora koji se dodaje u kontekstu istraživanja o prihvatanju i korištenju generativne umjetne inteligencije, no istraživanje koje su proveli Foroughi i sur. (2023) pokazalo je kako se navedeni moderator nije pokazao značajnim moderatorom u niti jednom odnosu.

Baptista i Oliveira (2015) za istraživanje u kontekstu mobilnog bankarstva UTAUT2 proširuju kulturološkim moderatorima, njih čak pet – individualizam/kolektivizam, izbjegavanje neizvjesnosti, dugoročna/kratkoročna orijentacija, muževnost/ženstvenost, distanca moći. Navedeni kulturološki moderatori, osim muževnosti/ženstvenosti, značajno su utjecali na ponašanje korisnika te su pomogli značajno povećati objašnjenje varijance modela s 45,9 % na 58,7 %.

Kada je riječ o korištenju novijih tehnologija, Nordhoff i sur. (2020) navode kako je u slučaju njihova istraživanja u kontekstu autonomnih vozila veća šansa da će ih prihvatiti muškarci nego

li žene, ali i da je veća šansa prihvaćanja kod mlađih osoba nego li kod starijih osoba. Korkmaz i sur. (2022) proveli su istraživanje na srodnoj temi, na prihvaćanju i korištenju sustava autonomnih vozila u javnom prijevozu te zaključili kako u samo dva slučaja moderatori imaju utjecaj na namjeru ponašanja, a riječ je o hedonističkoj motivaciji kod ženskog spola te povjerenje i sigurnost kod muškog spola.

Schmitz i sur. (2022) navode kako je njihovo istraživanje pokazalo da hedonistička motivacija i percipirana sigurnost značajno predviđaju namjeru korištenja kod mlađih skupina, dok su očekivani naponi i očekivana izvedba prediktori namjere ponašanja kod starijih skupina. Isto istraživanje, koje se provodilo u kontekstu namjere korištenja virtualnih medicinskih pregleda, utvrdilo je da su muškarcima hedonistička motivacija i očekivanja izvedbe bitni čimbenici koji utječu na namjeru korištenja navedene usluge, dok je kod žena to hedonistička motivacija i percipirana sigurnost. Vinerean i sur. (2022) navodi kako mlađe generacije, konkretno generacija Z, nema problema s povjerenjem u aplikacije m-trgovine te to njihovo istaknuto povjerenje vodi do jasne namjere da nastave ovaj oblik kupnje.

Istraživanje koje su proveli Nikolopoulou i sur. (2020) pokazuje kako istraživanje provedeno na studentskoj populaciji ne pokazuje kako dob, spol i iskustvo imaju utjecaj na prediktore namjere ponašanja ili stvarno korištenje jer je ispitani uzorak previše homogen. Lee i sur. (2024) također navode kako su koristili različite demografske karakteristike (dob, spol, obrazovanje) kao moderatore, ali niti jedan odnos na kojima su testirane nije se pokazao značajnim. Vrlo sličan zaključak pokazalo je istraživanje koje je proveo Strzelecki (2023), gdje je na studentskoj populaciji ispitivao spremnost na korištenje i prihvaćanje ChatGPT-a te se pokazalo da niti godina studija niti spol nisu značajni moderatori na nijedan odnos među varijablama. Grassini i sur. (2024) također su istraživali prihvaćanje i korištenje ChatGPT-a među studentima te je i njihovo istraživanje pokazalo kako dob, spol i iskustvo ne moderiraju niti posreduju nijedan od odnosa na kojima su testirane, a isti zaključak navode i Lavidas i sur. (2024) koji su ispitivali prihvaćanje i korištenje aplikacija umjetne inteligencije među studentima društvenih i humanističkih znanosti. Istraživanje koje su proveli Biloš i Budimir (2024) isključilo je moderatore iz modela jer se niti jedan od tri standardna moderatora (dob, spol i iskustvo) nisu pokazali značajnima niti na jednom odnosu među varijablama.

Cabrera-Sanchez i sur. (2021) strah od tehnologije i korisničko povjerenje koriste kao moderator, ali i kao medijator u istraživanju o prihvaćanju u korištenju umjetne inteligencije. No, mnogi drugi radovi koriste neke od standardnih moderatora kao medijatore ili kontrolne varijable (Sugumar & Chandra, 2021; Cao i sur., 2023; Choudhary i sur., 2024). Ipak, u

nastavku ovog rada korišteni su standardni moderatori za UTAUT2 – dob, spol i iskustvo (Venkatesh i sur., 2012).

4.3. Analiza prethodnih istraživanja

Venkatesh i sur. (2012) navode kako su čestice za četiri osnovna konstrukta preuzete iz rada u kojem je izvorno predstavljen UTAUT (Venkatesh i sur., 2003). Ipak, kako je ranije navedeno, postoje manje modifikacije spomenutih čestica koje su Venkatesh i sur. (2012) usavršili unutar UTAUT2 modela. Tri konstrukta koja su Venkatesh i sur. (2012) nadodali u UTAUT2 su hedonistička motivacija (HM), cijena (PV) i navika (HT). Autori modela navode kako su čestice za konstrukt navike (HT) preuzeta od Limayema i Hirta (2003), čestice za hedonističku motivaciju (HM) preuzete su iz istraživanja koje su proveli Kim i sur. (2005), a čestice za vrijednost cijene (PV) prilagođene su po uzoru na Dodds i sur. (1991). Unutar ovog doktorskog rada čestice iz izvornog UTAUT2 modela koristit će se u gotovo istom obliku, no prilagođene kontekstu istraživanja. Također, istraživački instrument će biti na engleskom i na hrvatskom jeziku. Osim navedenih konstrukta koji su prethodno nabrojani, ovaj rad bit će proširen konstruktima osobne inovativnosti (PI) i povjerenja (TR). Nadalje, čestice koje neće biti preuzete po uzoru na izvorni UTAUT2 su čestice za konstrukt zavisne varijable stvarnog korisničkog ponašanja (USE), čestice za zavisnu varijablu bit će modificirane po uzoru na Venkatesh i sur. (2008).

U prethodno napravljenom pregledu literature, od radova u kojima je bilo jasno navedeno kakvu su ljestvicu koristili, njih 39 koristi Likertovu ljestvicu od 1 do 5, a 22 rada koriste Likertovu ljestvicu od 1 do 7 te jedan rad koji koristi ljestvicu od 0 do 100. Dakle, od ukupnog broja radova koji u radu navode korištenu ljestvicu, njih 63 % koristi Likertovu ljestvicu od 1 do 5, a njih 35 % koristi Likertovu ljestvicu od 1 do 7. Ipak, u navedeni popis nisu ubrojani radovi Venkatesha i sur. (2003) i Venkatesha i sur. (2012) koji su temelj ovoga modela, a koji također koriste Likertovu ljestvicu od 1 do 7. S obzirom na to da je ranije navedeno kako će osnovni okvir konceptualnog modela biti preuzet iz Venkatesh i sur. (2012), isto će biti učinjeno i s mjernom ljestvicom. Osim spomenute Likertove ljestvice i čestica konstrukta, radi usporedivosti, po uzoru na istraživanje Venkatesh i sur. (2012) preuzet će se i mjerenja za moderatore pa će ponuđene opcije kod spola biti muškarci i žene, dob će biti mjerena u godinama, a iskustvo će se mjeriti u mjesecima korištenja alata generativne umjetne inteligencije.

Venkatesh i sur. (2012) navode kako su za testiranje modela koristili PLS metodu (PLS; engl. – *Partial Least Square* / hrv. Metoda parcijalnih najmanjih kvadrata). PLS ili PLS-SEM je metoda strukturalnog modeliranja korištena kod modela s mnoštvom latentnih varijabli i indikatora, ali je i sjajna metoda za otkrivanje odnosa među varijablama, a za razliku od CB-SEM-a ne zahtijeva jako velike uzorke. Od ukupno 68 analiziranih radova u pregledu literature, čak 48 (71 %) njih koristi PLS, odnosno PLS-SEM za testiranje modela.

Prosječna ekstrahirana varijanca (AVE – engl. *Average variance extracted*) druga je iznimno važna statistička mjera koja se koristi u 56 od ukupno 68 analiziranih radova (82 %). Riječ je o statističkoj mjeri koja se u pravilu koristi kod strukturalnog modeliranja kako bi se procijenila valjanost latentnih varijabli. Nadalje, u 30 od ukupno analiziranih 68 radova (44 %) navodi se korištenje korištenje faktora inflacije varijance (VIF – engl. *Variance inflation factor*). VIF mjeri koliko je varijanca procijenjenog koeficijenta regresije povećana zbog prisutnosti multikolinearnosti, a što nije rijetka pojava kada su neovisne varijable u modelu međusobno povezane. Biloš i Budimir (2024) u istraživanju navode postojanje previsoke razine multikolinearnosti između namjere ponašanja i stvarnog korisničkog ponašanja.

Tablica 13. Pregled radova koji se bave generativnom umjetnom inteligencijom

Izvor	Kontekst	Objašnjenja varijance	Uzorak
Al-Emran (2023)	Chatbotovi pogonjeni umjetnom inteligencijom. (Malezija)	USE – 57 %	Malezija, 447 ispitanika, studentska populacija
Biloš & Budimir (2024)	ChatGPT (Gen AI)	BI – 65 %	Hrvatska, 694 ispitanika, Generacija Z
Budhathoki i sur. (2024)	ChatGPT (Gen AI)	BI – 69,6 % USE – 74,2 %	Nepal i UK, 470 ispitanika, Nepal – 239, UK – 226
Foroughi i sur. (2024)	ChatGPT (Gen AI)	BI – 52,7 %	Malezija, 406 ispitanika, studentska populacija

Garcia de Blanes Sebastian i sur. (2022)	Virtualni glasovni asistenti pogonjeni umjetnom inteligencijom	BI – 89,8 %	Španjolska, 306 ispitanika
Goyal i sur. (2023)	ChatGPT (Gen AI)	USE – 56 %	Indija, 149 ispitanika, studentska populacija
Grassini i sur. (2024)	ChatGPT (Gen AI)	BI – 86,6 % USE – 65,6 %	Norveška, 104 ispitanika, studentska populacija
Habibi i sur. (2023)	ChatGPT (Gen AI)	BI – 40 % USE – 41,6 %	Indonezija, 1117 ispitanika, studentska populacija
Lai i sur. (2024)	ChatGPT (Gen AI)	Nema podataka	Hong Kong, 483 ispitanika, studentska populacija
Lavidas i sur. (2024)	ChatGPT (Gen AI)	BI – 75 % USE – 67 %	Grčka, 197 ispitanika, studentska populacija
Lee i sur. (2024)	ChatGPT (Gen AI)	BI – 62 % (ATT – 75 %)	SAD, 1004 ispitanika
Maican i sur. (2023)	Gen AI za slike	BI – 69 %	Rumunjska, 403 ispitanika, studentska populacija
Mehedi Hasan Emon i sur. (2023)	ChatGPT (Gen AI)	BI – 69,8 % USE – 80,6 %	Bangladeš, 350 ispitanika, zaposlene osobe
Menon & Shilpa (2023)	ChatGPT (Gen AI)	Nema podataka	Indija, 32 ispitanika
Mohd Rahim i sur. (2022)	Chatbot	BI – 66,9 % USE – 62,2%	Malezija, 334 ispitanika, studentska populacija

Salifu i sur. (2024)	ChatGPT (Gen AI)	BI – 74 % USE – 52 %	Gana, 306 ispitanika, studentska populacija
Sallam i sur. (2023)	ChatGPT (Gen AI)	ATT – 69,3 % USE – 72 %	Jordan, 458 ispitanika, studentska populacija
Sinaga i sur. (2024)	ChatGPT (Gen AI)	Nema podataka	Indonezija, 339 ispitanika, studentska populacija
Sobaih i sur. (2024)	ChatGPT (Gen AI)	BI – 30,3 % USE – 74 %	Saudijska Arabija, 520 ispitanika, studentska populacija
Strzelecki (2023)	ChatGPT (Gen AI)	BI – 73,4 %	Poljska, 534 ispitanika, studentska populacija
Sugumar & Chandra (2021)	Chatbot za financijske usluge	BI – 65 %	Indija, SAD, Singapur, 250 ispitanika
Vimalkumar i sur. (2021)	Virtualni glasovni asistenti	BI – 53 % USE – 14 %	Indija, 252 ispitanika

Izvor: izrada autora

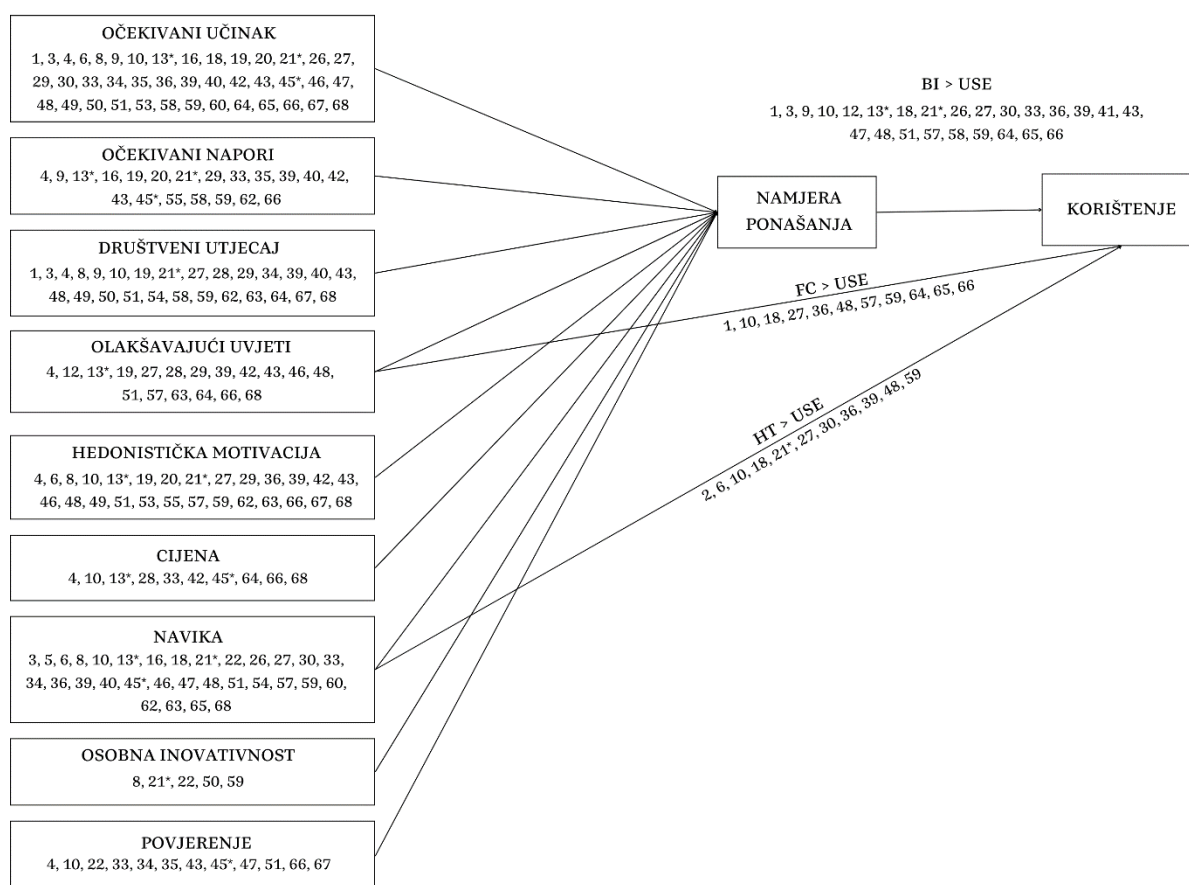
Od ukupno 22 rada koji su uzeti u obzir kao istraživanja koja istražuju korištenje i prihvaćanje generativne umjetne inteligencije, njih 16 (73 %) istražuje prihvaćanje i korištenje ChatGPT-a kao najprepoznatljivijeg alata generativne umjetne inteligencije. Vrlo je znakovito i kako 15 od ukupno 22 izdvojena rada (68 %) istraživanje provodi na studentskoj populaciji. Prosječan broj ispitanika u navedenim radovima je 416 ispitanika. Ogroman prostor za buduća istraživanja leži u činjenici da niti jedan od 22 analizirana rada ne istražuje alat generativne umjetne inteligencije na općoj populaciji s moderatorima uključenima u model. Koliko su moderatori u središtu UTAUT2 modela govori i činjenica da sve hipoteze koje su navedene u izvornim istraživanjima naglašavaju utjecaj moderatora kao srž istraživanja (Venkatesh i sur., 2003; Venkatesh i sur., 2012).

4.3.1. Odnosi u prethodnim istraživanjima

Za potrebe pregleda literature i usporedbe rezultata modela, pregledano je 68 prethodnih istraživanja na ovu i povezane teme, a u nastavku slika 13. i slika 14. prikazuju potvrđene i nepotvrđene utjecaje među varijablama. Brojke koje se nalaze ispod naziva konstrukta označavaju redni broj rada iz tablice 9.

Iz slike 13. vidljivo je kako je u analiziranim radovima konstrukt očekivanog učinka (PE) najčešće potvrđeni prediktor namjere ponašanja (BI). Nadalje, zanimljivo je za uočiti kako je utjecaj konstrukta očekivanih napora (EE) na namjeru ponašanja (BI) češće nepotvrđen nego potvrđen u pregledanim istraživanjima, a isto vrijedi i za olakšavajuće uvjete (FC). Navika (HT), društveni utjecaj (SI) i hedonistička motivacija (HM) češće su potvrđeni nego nepotvrđeni, a istraživanja koja obuhvaćaju osobnu inovativnost (PI) i povjerenje (TR) su prikazana u skromnijem obujmu, ali navedeni konstrukti u modelu češće su potvrđeni nego li nepotvrđeni.

Nadalje, za sva tri konstrukta koja utječu na zavisnu varijablu korištenja (USE), a to su namjera ponašanja (BI), olakšavajući uvjeti (FC) i navika (HT), vidljivo je kako je njihov utjecaj na varijablu korištenja u većoj mjeri potvrđen.



Slika 13. Pregled literature – prikaz utjecaja među varijablama koji su potvrđeni

Pri analizi prihvatanja i korištenja tehnologija kroz UTAUT2 model, uočavaju se određeni obrasci koji ukazuju na specifične faktore prihvatanja pojedinih tehnologija. Primjerice, kod glasovnih asistenata i *chatbotova*, značajnu ulogu ima hedonistička motivacija (HM) jer korisnici često koriste takve tehnologije iz znatiželje i zabave. Kod tehnologija za pametnu kuću, olakšavajući uvjeti (FC) i navika (HT) imaju snažan utjecaj jer se ove tehnologije integriraju u svakodnevne rutine korisnika. S druge strane, kod softverskih rješenja i AI platformi u poslovnom kontekstu, očekivani učinak (PE) nameće se kao najvažniji faktor prihvatanja takvih tehnologija. Kada su u središtu istraživanja tehnologije koje se direktno ili indirektno dotiču osobnih financija, faktori poput sigurnosti i povjerenja, ali i percipirana razina rizika igraju snažnu ulogu u prihvatanju i korištenju takve tehnologije.

Konstrukt očekivanog učinka (PE), osim što je najčešće potvrđeni prediktor u ovoj analizi, najčešće je i najsnažniji prediktor namjere ponašanja, posebno kod namjere korištenja umjetne inteligencije (Gansser & Reich, 2021; Cabrera-Sanchez i sur., 2021; Grassini i sur., 2024; Lavidas i sur., 2024; Sobaih i sur., 2024). U skladu je to i sa zaključcima koje navode Venkatesh i sur. (2003) u izvornom UTAUT-u, a koji kažu kako je očekivani učinak najsnažniji prediktor namjere korištenja nove tehnologije. Značajke alata generativne umjetne inteligencije su da štede vrijeme, odnosno povećavaju vremensku efikasnost pojedinca, poboljšavaju učinkovitost i povećavaju mogućnosti rješavanja raznih zadataka (Budhathoki i sur., 2024; Bin-Nashwan i sur., 2023; Chiu i sur., 2023). Ukoliko su ispitanici svjesni da tehnologija može povećati njihovu učinkovitost, to će utjecati na njihovu namjeru korištenja navedene tehnologije (Venkatesh i sur., 2003; Budhathoki i sur., 2024). Prednosti korištenja alata generativne umjetne inteligencije u vidu povećanja produktivnosti vidljivi su u mnogim sferama života, od obveza u svakodnevnom privatnom životu do obrazovanja, ali i raznih poslovnih obveza (Grassini i sur., 2024; Gansser & Reich, 2021; Maican i sur., 2023; Lavidas i sur., 2024; Mohd-Rahim i sur., 2022; Budhathoki i sur., 2024).

Očekivani napori (EE) definiraju se kao stupanj lakoće ili jednostavnosti povezan s prihvatanjem i korištenjem tehnologije (Venkatesh i sur., 2003). Premda teorija kaže kako jednostavnost rukovanja tehnologijom povećava namjere pojedinaca da koriste tehnologiju i da ih demotivira ikakva tehnologija koju je malo kompleksnije za koristiti (Cimperman i sur., 2013; Strzelecki, 2023; Budhathoki i sur., 2024), pregled literature sugerira kako očekivani napori često nemaju učinak na namjeru ponašanja. Menon i Shilpa (2023) navode kako mladi korisnici alate poput ChatGPT-a doživljavaju jednostavnima i intuitivnima te da to povećava njihovu namjeru korištenja takvih alata, a potvrđuju to i druga istraživanja o mlađoj populaciji

(Sobaih i sur., 2024; Strzelecki, 2023; Lai i sur., 2024). Maican i sur. (2023) navode zanimljive zaključke kako kod AI alata za generiranje slika očekivani naponi (EE) imaju snažniji utjecaj na namjeru ponašanja (BI) kod manje kreativnih pojedinaca, ali i kod onih koji imaju naprednije znanje engleskog jezika (što vrijedi kod alata koji ne funkcioniraju na više jezika, nego samo na engleskom jeziku, kao npr. Midjourney).

Društveni utjecaj (SI) opisuje se kao stupanj u kojem pojedinac osjeća važnost koju drugi pridaju korištenju određene tehnologije ili koliko važni ljudi iz njegovog okruženja smatraju da bi i on trebao koristiti tu tehnologiju (Venkatesh i sur., 2003; Venkatesh i sur., 2012). Društveni utjecaj očituje se u tome da pojedinci promatraju stavove, mišljenja i ponašanja ljudi oko sebe u vezi s integracijom određene tehnologije, primjerice generativne umjetne inteligencije u privatnom i poslovnom životu (Venkatesh, 2022). Istraživanja potvrđuju kako utjecaj okoline utječe na pojedince da prihvate alate generativne umjetne inteligencije (Sobaih i sur., 2024; Xian, 2021; Budhathoki i sur., 2024), premda se ovaj konstrukt češće pokazuje važnim kod telekomunikacijskih tehnologija gdje pojedinci žele koristiti iste komunikacijske kanale kao i njihovi prijatelji i obitelj (Ameri i sur., 2020; Nikolopoulou i sur., 2020).

Konstrukt olakšavajućih uvjeta (FC) odnosi se na percepciju postojećih infrastrukturnih resursa i dostupne pomoći potrebne za korištenje tehnologije (Venkatesh i sur., 2003). U UTAUT2 modelu olakšavajući uvjeti prediktor su i namjere ponašanja, ali i stvarnog korištenja (Venkatesh i sur., 2012). U kontekstu korištenja alata generativne umjetne inteligencije, olakšavajući uvjeti odnose se na postojeću tehničku infrastrukturu koja je nužna za korištenje takve tehnologije (Mehedi Hasan Emon i sur., 2023), ali odnosi se i na druge resurse poput znanja, obuke, dostupnosti virtualnih instrukcija, uputstava te mogućnosti da se zatraži pomoć od digitalno pismenije osobe iz okruženja (Xian, 2021; Faqid & Jaradat, 2021). Neka od istraživanja na studentskoj populaciji pokazuju kako olakšavajući uvjeti pozitivno utječu na namjeru ponašanja (Habibi i sur., 2023; Salifu i sur., 2023; Chatterjee & Bhattacharjee, 2020; Sinaga i sur., 2024), ali također i na stvarno korištenje (Habibi i sur., 2023; Salifu i sur., 2023; Strzelecki, 2023; Lavidas i sur., 2024). Ipak, analiza pregleda literature ukazuje na to da olakšavajući uvjeti u manjoj mjeri imaju značajan utjecaj na namjeru ponašanja, ali često imaju značajan utjecaj na stvarno korištenje.

Izvorni UTAUT sadrži konstrukt očekivanog učinka (PE) koji je čimbenik vanjske motivacije i dominantno se usredotočuje na korisnost određene tehnologije koja je u središtu istraživanja (Venkatesh i sur., 2003). U UTAUT2 dodan je konstrukt hedonističke motivacije (HM) koja se oslanja na unutarnji čimbenik motivacije za korištenje određene tehnologije (Venkatesh i sur.,

2012). Veća razina hedonističke motivacije u korištenju neke tehnologije poput umjetne inteligencije povećava namjeru korisnika da koristi takvu tehnologiju (Gansser & Reich, 2021). Mogućnosti alata generativne umjetne inteligencije su različiti, ali kod glasovnih asistenata hedonistička motivacija najsnažniji je prediktor namjere, snažnija čak i od očekivanog učinka kao vanjskog motivatora (Vimalkumar i sur., 2021). No, i istraživanja na drugim alatima poput ChatGPT-a pokazuju kako interakcija na relaciji čovjek – umjetna inteligencija ljudima stvara užitek i zabavu te je zbog toga bitan prediktor namjere ponašanja (Strzelecki, 2023; Habibi i sur., 2023; Mehedi Hasan Emon i sur., 2023; Sinaga i sur., 2024).

Cijena ili vrijednost cijene (PV) još je jedan od konstrukta koji je novitet u UTAUT2 modelu u odnosu na izvorni UTAUT, a riječ je o kontekstualnom faktoru koji opisuje percipirani odnos troška i vrijednosti koju tehnologija isporučuje (Venkatesh i sur., 2012; Cabrera – Sanchez i sur., 2021). Utjecaj cijene pozitivan je kada korisnik prednosti korištenja tehnologije smatra važnijima od izdvojenih novčanih sredstava (Baptista & Oliveira, 2015). Ukoliko je cijena visoka, ona može imati negativan utjecaj na namjeru za korištenjem tehnologije (Brown & Venkatesh, 2005). Kako je vidljivo iz analize, cijena relativno često nije prediktor namjere korištenja, a značajna je uglavnom kada je riječ o digitalnim proizvodima i uslugama koji za relativno nisku cijenu pružaju i isporučuju dobru vrijednost, kao što su primjerice usluge mobilnog bankarstva i digitalnog novčanika (Merhi i sur., 2019; Hassan i sur., 2023; Kilani i sur., 2023), internetske platforme za edukaciju (Tseng i sur., 2022; Meet i sur., 2022) ili mobilne videoigre (Baabdullah (2018)). Ipak, većina promatranih istraživanja o alatima generativne umjetne inteligencije nastala je u periodu kada je jako mali udio ispitanika koristio plaćene verzije takvih alata i taj se postotak kontinuirano povećava (Law, 2024). S obzirom na sve veće mogućnosti koje takvi alati pružaju, u budućnosti će biti zanimljivo promatrati odnos percipiranog odnosa cijene i isporučene vrijednosti.

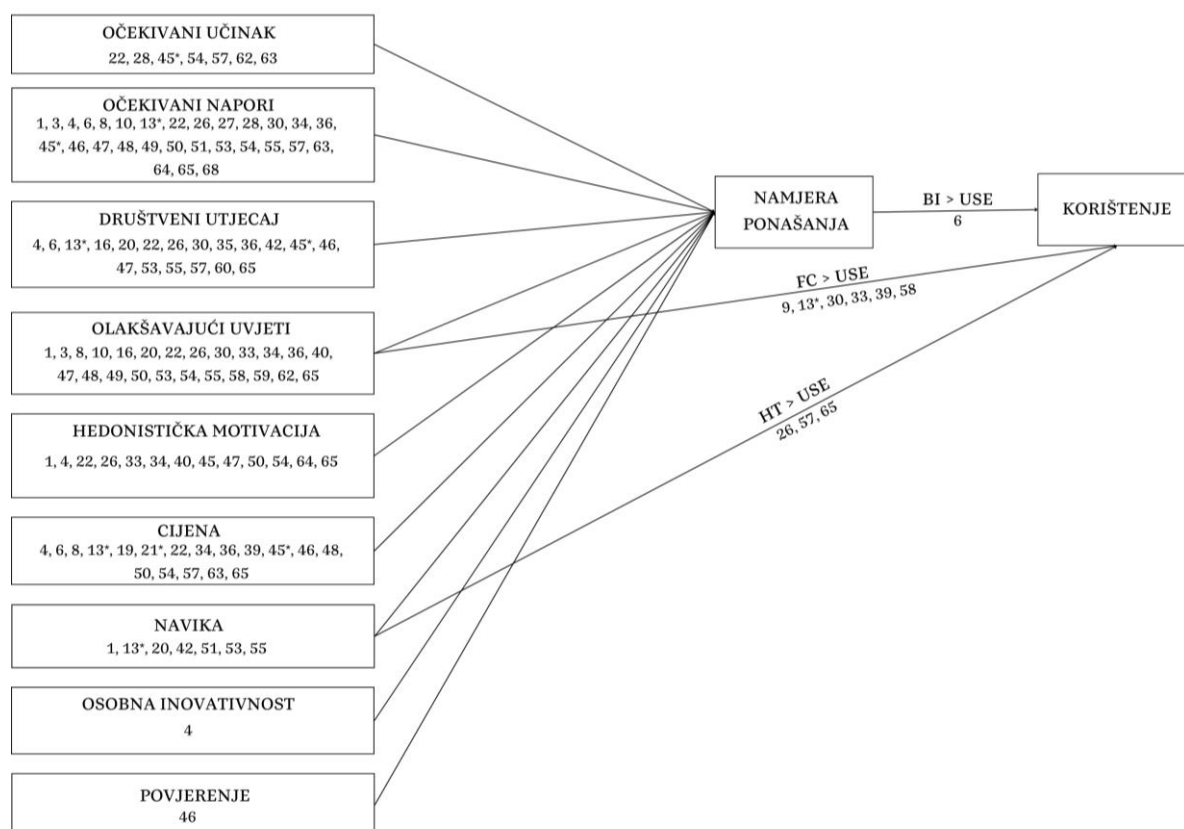
Utjecaj navike (HT) koji se definira kao stupanj u kojemu ljudi automatski izvode ponašanja jer su iskustveno tako naučili (Limayem i sur., 2007), utječe i na namjeru ponašanja, ali i na stvarno ponašanje (Venkatesh i sur., 2012). U skladu sa zaključcima iz planirane teorije ponašanja (Ajzen, 1991), pretpostavlja se kako ponavljajuće ponašanje može uzrokovati ukorijenjene navike koje su okidač za namjeru ponašanja, ali u konačnici i stvarnog ponašanja (Venkatesh i sur., 2012; Gansser & Reich, 2021). Iz priložene analize pregleda literature, vidljivo je kako veći dio istraživanja potvrđuje tu pretpostavku. U određenim istraživanjima kao što su primjerice ona o korištenju mobilnih aplikacija, navika je nerijetko i najsnažniji prediktor namjere ponašanja (Hew i sur., 2015; Nikolopoulou i sur., 2020; Merhi i sur., 2019).

Navika može biti najsnažniji prediktor namjere korištenja i kod uređaja koji koriste umjetnu inteligenciju, poput autonomnih vozila ili zdravstvenih uređaja (Gansser & Reich, 2021; Garcia de Blanes Sebastian i sur., 2022; Baudier i sur., 2020). Istraživanja pokazuju kako je i kod korištenja alata generativne umjetne inteligencije navika jedan bitan prediktor namjere ponašanja (Strzelecki, 2023; Salifu i sur., 2024; Sinaga i sur., 2024; Lavidas i sur., 2024; Grassini i sur., 2024; Sugumar & Chandra, 2021; Maican i sur., 2023; Habibi i sur., 2023). Maican i sur. (2023) navode kako muškarci koji su razvili naviku korištenja alata generativne umjetne inteligencije imaju veću vjerojatnost da će tu tehnologiju koristiti i u budućnosti. Navika također ima direktni utjecaj na stvarno korištenje (Venkatesh i sur., 2012), a što se potvrdilo i u drugim istraživanjima na raznim tehnologijama, uključujući i one koje pogoni umjetna inteligencija (Gansser & Reich, 2021; Al-Emran i sur., 2023; Cabrera-Sanchez i sur., 2021, Habibi i sur., 2023; Strzelecki, 2023; Lavidas i sur., 2024) Ipak, određena istraživanja ukazuju na činjenicu da navika utječe na namjeru korištenja, ali ne i stvarno korištenje, što je također zanimljiv ishod istraživanja (Sinaga i sur., 2023; Grassini i sur., 2024; Thusi i Maduku, 2020).

Konstrukt osobne inovativnosti (PI), koji nije sastavni dio UTAUT2 modela već nadodani konstrukt u konceptualnom modelu ovog istraživanja, korišten je tek nekoliko puta u radovima unutar analizirane literature. Agarwal i Prasad (1998) definiraju ga kao stupanj u kojem je pojedinac spreman na isprobavanje novih tehnologija. Premda u istraživanjima on najčešće nije osobito snažan prediktor, ipak je statistički značajan i pozitivan prediktor namjere ponašanja (Gansser i sur., 2021; Garcia de Blanes Sebastian i sur., 2022; Strzelecki i sur., 2023).

Povjerenje je također novopridruženi konstrukt na originalni UTAUT2, a on predstavlja stupanj sigurnosti i povjerenja koju osoba stavlja u korištenje proizvoda ili usluge (Kim i sur., 2017). S obzirom na to da istraživanje iz Republike Hrvatske ukazuje na činjenicu da je povjerenje u umjetnu inteligenciju vrlo nisko (Kopal i sur., 2024), konstrukt povjerenja dodan je da istraži koliko navedeno povjerenje ima utjecaj na namjeru ponašanja. Kod prihvatanja tehnologije koje sa sobom nose visoki rizik, poput autonomnih vozila ili digitalnog novčanika, povjerenje može imati ključnu ulogu (Korkmaz i sur., 2022; Kilani i sur., 2023). Ipak, kod tehnologija koje sa sobom nose određeni rizik, poput alata generativne umjetne inteligencije i srodnih tehnologija, povjerenje i drugi sigurnosti faktori imaju manju važnost u modelu, ali i dalje utječu na namjeru ponašanja (Mehedi Hasan Emon i sur., 2023; Lai i sur., 2024; Salifu i sur., 2024; Cabrera-Sanchez i sur., 2021, Garcia de Blanes Sebastian i sur., 2022; Mohd Rahim i sur., 2022; Vimalkumar i sur., 2021).

Namjera ponašanja prethodi stvarnom korištenju, a temelji se na teoriji racionalnog djelovanja (TRA) (Fishbein i sur., 1975). Sheppard i sur., (1988) u svojoj meta-analizi na 87 zasebnih istraživanja ispituju odnos između navedenih varijabli i potvrđuju pretpostavku kako postoji snažna povezanost između namjere ponašanja i stvarnog korištenja. S obzirom na to da Venkatesh i sur. (2012) ne navode čestice za konstrukt stvarnog korištenja, autori koji su integrirali UTAUT2 koriste različite verzije ove čestice i rezultati se uvelike razlikuju. Veliki broj radova iz pregleda literature uopće ne koristi konstrukt stvarnog ponašanja, ali radovi koji koriste gotovo svi potvrđuju utjecaj namjere ponašanja na stvarno korištenje. Utjecaj namjere ponašanja na stvarno korištenje potvrdila su i istraživanja u čijem su središtu alati generativne umjetne inteligencije (Habibi i sur., 2023; Salifu i sur., 2024; Strzelecki, 2023; Budhathoki i sur., 2024; Grassini i sur., 2024; Mehedi Hasan Emon i sur., 2023; Sinaga i sur., 2024; Sobaih i sur., 2024; Cabrera-Sanchez, 2021; Chatterjee i Bhattacharjee, 2020; Mohd Rahim i sur., 2022).



Slika 14. Pregled literature – prikaz utjecaja među varijablama koji nisu potvrđeni

Kako je i vidljivo iz slike 14., u istraživanjima koja su u središtu ovog pregleda literature spominju se i brojni drugi konstrukti koji nisu predmet obrade u slici 13. i slici 14. Navedeni

pregled odnosa među varijablama usmjerava se samo na konstrukte koji su dio konceptualnog modela ovoga rada.

Iz navedenih je grafikona vidljivo kako očekivani naponi (EE), olakšavajući uvjeti (FC) i cijena (PV) češće ne utječu na namjeru ponašanja nego li što utječu. Za konstrukt očekivanih napora (EE), Sobaih i sur. (2024) te Strzelecki (2023) navode kako ispitanici ChatGPT-a, posebno mladi ispitanici, navedeni alat generativne umjetne inteligencije doživljavaju kao lako razumljivim te jednostavnim za korištenje i komunikaciju i zbog toga bi njihova namjera korištenja trebala biti pozitivna. Thusi i Maduku (2020) navode kako mlađe generacije većinu tehnologije doživljavaju jednostavnom za korištenje te zbog toga navode kako ne ulažu poseban napor u svladavanje tehnologije. Ameri i sur. (2019) navode da faktor očekivanih napora korisnici ne doživljavaju bitnim ukoliko je korištenje određene tehnologije ili aplikacije prirodno, jasno i razumljivo. Različita su tumačenja autora oko toga zašto utjecaj očekivanih napora na namjeru ponašanja funkcionira ili ne funkcionira, ali iz navedenih je grafikona vidljivo kako on češće ne funkcionira. Bitno je i napomenuti kako spomenuti utjecaj ovisi i o heterogenosti istraživane populacije (Nikolopoulou i sur. 2020), ali naravno i o tehnologiji koja je u središtu istraživanja. Mehedi Hasan Emon i sur. (2023) navode kako kod korištenja ChatGPT-a olakšavajući uvjeti pozitivno utječu na namjeru korištenja i to kao najsnažniji prediktor, ali da očekivani naponi čak negativno utječu na namjeru korištenja. Maican i sur. (2023) navode kako je negativni utjecaj očekivanih napora na namjeru ponašanja snažniji kod muškaraca.

Shaw i Sergueeva (2019) potvrđuju tvrdnje Venkatesha i sur. (2003) koji navode kako ispitanici mogu poistovjetiti olakšavajuće uvjete (FC) i očekivane napore (EE). Osim što je alate generativne umjetne inteligencije relativno jednostavno koristiti, za njihovo korištenje nisu potrebni osobito značajni resursi niti snažni uređaji (Stzelecki, 2024). Većina modernih tehnologija namijenjenih široj publici, posebno softverska rješenja, danas nudi vrlo intuitivna sučelja koja rijetko zahtijevaju pomoć korisničkih službi, servisnih centara ili pomoć kolega koji su spretniji s tehnologijom (Shaw i Sergueeva, 2019). Ukoliko tehnološki problemi ne predstavljaju izazovan zadatak za ispitanike, to se odražava i na utjecaj olakšavajućih uvjeta na namjeru ponašanja (Schmitz i sur., 2022). Ukoliko korisnici alate generativne umjetne inteligencije percipiraju jednostavnima za korištenje, to eliminira potrebu za podržavajućom infrastrukturu i to objašnjava zašto olakšavajući uvjeti nemaju utjecaj na namjeru ponašanja u istraživanja o prihvaćanju tehnologija koje su jednostavne za korištenje (Mohd Rahim i sur., 2022). Ipak, pregled literature ukazuje na činjenicu da je više istraživanja gdje olakšavajući

uvjeti imaju pozitivan utjecaj na stvarno korištenje, nego onih gdje nemaju, ali i da nije neuobičajeno da unutar istog istraživanja olakšavajući uvjeti utječu na korištenje, ali ne i na namjeru korištenja (Ain, 2025; Cabrera-Sanchez i sur., 2021; Nikolopoulou i sur., 2020; Strzelecki (2023); Lavidas i sur., 2024; Sinaga i sur., 2024).

Cijena (PV) je također jedan od faktora koji je u pregledu literature češće pokazao da nema pozitivan utjecaj na namjeru ponašanja. Baptista i Oliveira (2015) navode kako cijena nema značajan utjecaj na namjeru ponašanja ukoliko ispitanici tehnologiju ne plaćaju izravno. Isti zaključak navode i Garcia de Blanes Sebastian i sur. (2020). Thusi i Maduku (2020) u istraživanju o prihvaćanju i korištenju aplikacija za mobilno bankarstvo navode sličan zaključak jer banke ne naplaćuju samo korištenje aplikacije, premda naplaćuju svoje provizije kroz druge aktivnosti unutar aplikacije. Strzelecki (2023) navodi kako je on konstrukt cijene izbacio iz modela jer većina korisnika ChatGPT koristi besplatno.

Premda faktor društvenog utjecaja (SI) češće funkcionira nego li ne funkcionira, u određenim istraživanjima nije pokazao značajan utjecaj na namjeru ponašanja (BI) (Lavidas i sur., 2024; Schmitz i sur., 2022; Miladinović i Xiang, 2016; Kilani i sur., 2023; Kasparova, 2022). Grassini i sur. (2024) navode kako neki od mogućih razloga zašto društveni utjecaji (SI) ne utječu na namjeru ponašanja (BI) leže u činjenici da alate generativne umjetne inteligencije češće koriste visokoobrazovani, a oni su pod manjim utjecajem vanjskih društvenih pritisaka, ali i zbog toga što takvi alati još uvijek nemaju status sveopće prihvaćenog proizvoda koji bi svi trebali koristiti. Thusi i Maduku (2020) navode kako mlađe generacije manje brinu o mišljenjima svoje okoline, a Mohd Rahim i sur., (2022) dodaju kako generacija Z ne mari previše o mišljenjima, savjetima i preporukama drugih kada je u pitanju korištenje AI alata.

Neki od najčešćih nedostataka koje autori navode su ograničen uzorak ispitanika, uglavnom premalo prikupljenih ispitanika ili previše homogena skupina ispitanika što negativno utječe na generalizaciju podataka (Ain i sur., 2016; Garcia de Blanes Sebastian i sur., 2022). Osim toga, ograničenja kod radova mogu biti ograničenost na određeno geografsko područje ili na industriju što također smanjuje mogućnost generalizacije rezultata (Kilani i sur., 2023; Vinerean i sur., 2022).

Kao neke od najčešćih preporuka za buduća istraživanja navodi se provođenje longitudinalnih istraživanja u kojima bi se dugoročno pratilo prihvaćanje tehnologije kroz vrijeme te kako bi se kroz to identificirali faktori koji utječu na promjene u korisničkim navikama tijekom vremena (Choudhary i sur., 2024; Schmitz i sur., 2022; Baudier i sur., 2020; Korkmaz i sur., 2022).

4.3.2. Preliminarno istraživanje

Autor ovog doktorskog rada proveo je i preliminarno istraživanje pri čemu je istraživao korištenje specifičnog alata generativne umjetne inteligencije, ChatGPT-a, među pripadnicima generacije Z u Republici Hrvatskoj. Preliminarno istraživanje napravljeno je s ciljem testiranja utjecajnih čimbenika konceptualnog modela, ali i sa željom publiciranja rada u časopisu otvorenog pristupa, čemu je prethodio recenzenski postupak koji je rezultirao smjernicama za usavršavanje koncepta rada. U travnju 2024. godine objavljen je znanstveni rad pod nazivom „*Understanding the Adoption Dynamics of ChatGPT among Generation Z: Insights from a Modified UTAUT2 model*“ u časopisu *Journal of Theoretical and Applied Electronic Commerce Research* u koautorstvu s profesorom Bilošem (Biloš i Budimir, 2024).

Tijekom svibnja 2023. godine provodilo se kvantitativno istraživanje pri čemu su podatci prikupljeni CAWI metodom (engl. *Computer-assisted web interviewing*). Internetska anketa izrađena je u platformi Alchemer. Prilikom prikupljanja podataka korištena je metoda snježne grude pri čemu su studenti Sveučilišta Josipa Jurja Strossmayera u Osijeku po jasnim uputama širili anketu svojim poznanicima, pripadnicima generacije Z. Ukupno je prikupljeno 1159 ispitanika, no nakon filtracije podataka i eliminacije nepotpunih upitnika i diskvalifikaciju ispitanika koji nikada ranije nisu koristili ChatGPT, preostalo je 694 ispunjena upitnika koji su konačno analizirani.

U navedenom preliminarnom istraživanju korišten je UTAUT2 model koji je sadržavao očekivani učinak (PE), očekivane napore (EE), društveni utjecaj (SI), olakšavajuće uvjete (FC), hedonističku motivaciju (HM), cijenu (PV), naviku (HT) i osobnu inovativnost (PI), ali i moderatore dob, spol i iskustvo te zavisne varijable namjere ponašanja (BI) i stvarno korištenje (USE).

Iz publiciranog znanstvenog rada izostavljeni su moderatori i konstrukt stvarnog korištenja (USE). Tijekom analize podataka pokazalo se kako konstrukt stvarnog korištenja ima visoku multikolinearnost s konstruktom namjere ponašanja zbog čega je morala biti izbačena iz rada.

Za obradu podataka korišteni su statistički softverski paketi JASP i Jamovi, a analiza je uključivala nekoliko ključnih statističkih metoda:

- eksploratornu faktorsku analizu (EFA) – korištenu za provjeru strukture faktora i identifikaciju mogućih problema s korelacijom među varijablama

- konfirmatornu faktorsku analizu (CFA) – primjenjenu za validaciju mjernog modela i provjeru uklapanja podataka u predloženu strukturu faktora
- hijerarhijsku linearnu regresiju – korištenu za testiranje hipoteza i utvrđivanje značajnih prediktora namjere ponašanja.

Konačni model objasnio je 65 % varijance u namjeri korištenja ChatGPT-a, čime je potvrđena relativno visoka prediktivna snaga modificiranog UTUAT2 modela u kontekstu namjere usvajanja specifičnog alata generativne umjetne inteligencije. Međutim, određeni odnosi unutar modela nisu se pokazali značajnima što je dovelo do prilagodbi u interpretaciji rezultata.

Najznačajnijim prediktorom namjere ponašanja pokazala se navika (HT), što sugerira da se pripadnici generacije Z oslanjaju na ponavljajuće obrasce ponašanja u procesu usvajanja novih tehnologija. Očekivani učinak (PE) pokazao se kao drugi najsnažniji prediktor, potvrđujući da korisnici percipiraju alate generativne umjetne inteligencije kao sredstvo koje im može poboljšati produktivnost i olakšati obavljanje različitih zadataka. Hedonistička motivacija (HM) također je imala značajan utjecaj, što sugerira da korisnici ne koriste ChatGPT isključivo zbog njegove funkcionalnosti, već i zbog zabavnog aspekta interakcije s ovom tehnologijom. Društveni utjecaj (SI) također je imao pozitivan i značajan utjecaj, što je očekivano s obzirom na sklonost mladih ljudi da u donošenju odluka slijede stavove i ponašanja svoje društvene okoline. Osobna inovativnost (PI) varijabla je koju su Biloš i Budimir (2024) dodali na postojeći teorijski okvir UTAUT2 modela, a također se pokazala značajnim faktorom, sugerirajući da su korisnici koji su skloniji istraživanju novih tehnologija imali veću vjerojatnost prihvatanja alata generativne umjetne inteligencije, a zbog čega je preporučeno da se ona i dalje koristi u budućim istraživanjima.

Međutim, određeni aspekti modela nisu funkcionirali. Očekivani napor (EE) nije imao značajan utjecaj na namjeru korištenja ($p = 0,961$), što Biloš i Budimir (2024) objašnjavaju visokom razinom digitalne pismenosti generacije Z koja ne percipira ChatGPT kao tehnologiju koja zahtijeva poseban napor za korištenje. Olakšavajući uvjeti (FC) također nisu imali značajan učinak na namjeru ponašanja ($p = 0,652$). Ovaj rezultat može se objasniti činjenicom da generacija Z ima stalni pristup internetu i pametnim uređajima, što znači da ne doživljava tehničke prepreke pri korištenju ChatGPT-a, koji sam po sebi nije osobito zahtjevan softver za pokrenuti na gotovo svim uređajima. Naposljetku, vrijednost cijene ili cijena (PV) nije pokazala značajan utjecaj ($p = 0,561$), a Biloš i Budimir (2024) sugeriraju kako je ovaj rezultat posljedica činjenice da je ChatGPT u svom osnovnom obliku dostupan besplatno, čime cijena ne

predstavlja prepreku za korisnike. Ovaj nalaz je u skladu s istraživanjem Strzeleckog (2023), koji je iz istog razloga odlučio faktor cijene izbaciti iz modela.

Valjanost modela potvrđena je kroz eksploratornu (EFA) i konfirmatornu faktorsku analizu (CFA), koje su pokazale prihvatljiva faktorska opterećenja i međusobne korelacije varijabli. Hijerarhijska linearna regresija ukazala je na statistički značajne efekte pet od osam prediktora, čime je potvrđena teorijska osnova proširenog UTAUT2 modela u kontekstu generativne umjetne inteligencije.

Prilikom analize faktorske strukture modela, utvrđena je visoka korelacija između konstrukta namjere ponašanja (BI) i stvarnog korištenja (USE) (Pearsonov $r = 0,84$, $p < 0,001$). Zbog preklapanja među tim varijablama i potencijalnih problema s multikolinearnošću, odlučeno je da se varijabla stvarnog korištenja (USE) isključi iz daljnje analize. Čestice konstrukta stvarnog korištenja (USE) temeljile su se na česticama koje navode Nikolopoulou i sur. (2020), a navedene čestice su bile:

USE1 – Redovito koristim ChatGPT u svojim studijima.

USE2 – Korištenje ChatGPT-a je ugodno iskustvo.

USE3 – Trenutno koristim ChatGPT kao pomoći alat u svojim studijima.

USE4 – Provodim puno vremena koristeći ChatGPT.

Zbog visokih preklapanja s faktorom namjere ponašanja (BI) ovaj set čestica za mjerenje stvarnog korištenja predlaže se promijeniti u daljnjim istraživanjima. Biloš i Budimir (2024) objašnjavaju kako je ovaj rezultat vjerojatno posljedica samoprocjenske prirode korištenih čestica, pri čemu su ispitanici percipirali vlastitu namjeru ponašanja kao stvarno ponašanje. To upućuje na potrebu za korištenjem objektivnijih mjernih instrumenata u budućim istraživanjima kako bi se preciznije kvantificiralo stvarno korištenje ChatGPT-a i srodnih alata generativne umjetne inteligencije.

Neke od preporuka za buduća istraživanja, izuzev već spomenutog o promjeni čestica za zavisnu varijablu stvarnog ponašanja (USE), vežu se za proširenje uzorka na različite demografske skupine i detaljno proučavanja utjecaja moderatora na odnose unutar modela. S obzirom na to da se navedeno istraživanje usmjerilo samo na pripadnike generacije Z u Republici Hrvatskoj, navedena je skupina previše homogena da bi unutar nje moderator mogli funkcionirati. Zbog toga je ključno uzorak proširiti na različite demografske skupine. Nadalje, Biloš i Budimir (2024) dodaju kako bi se trebala razmotriti opcija integracije varijabli poput

povjerenja u umjetnu inteligenciju, osjećaja kontrole nad tehnologijom ili percipirane transparentnosti sustava umjetne inteligencije koje bi pružile dublji uvid u to kako korisnici donose odluke o prihvatanju generativne umjetne inteligencije.

5. USPOREDBA REPUBLIKE HRVATSKE I UJEDINJENOG KRALJEVSTVA PO SPREMNOSTI NA UMJETNU INTELIGENCIJU

Istraživanjem u okviru ovog doktorskog rada cilj je usporediti percepciju ispitanika iz Republike Hrvatske s još jednom zemljom kako bi se rezultati istraživanja o prihvaćanju i korištenju generativne umjetne inteligencije među hrvatskom populacijom mogli usporediti s rezultatima među ispitanicima koji su prema relevantnim dostupnim istraživanjima spremniji na umjetnu inteligenciju nego hrvatski ispitanici. Cilj takvog istraživanja bio bi odgonetnuti po čemu se prediktori u namjeri korištenja i korištenju generativne umjetne inteligencije razlikuju.

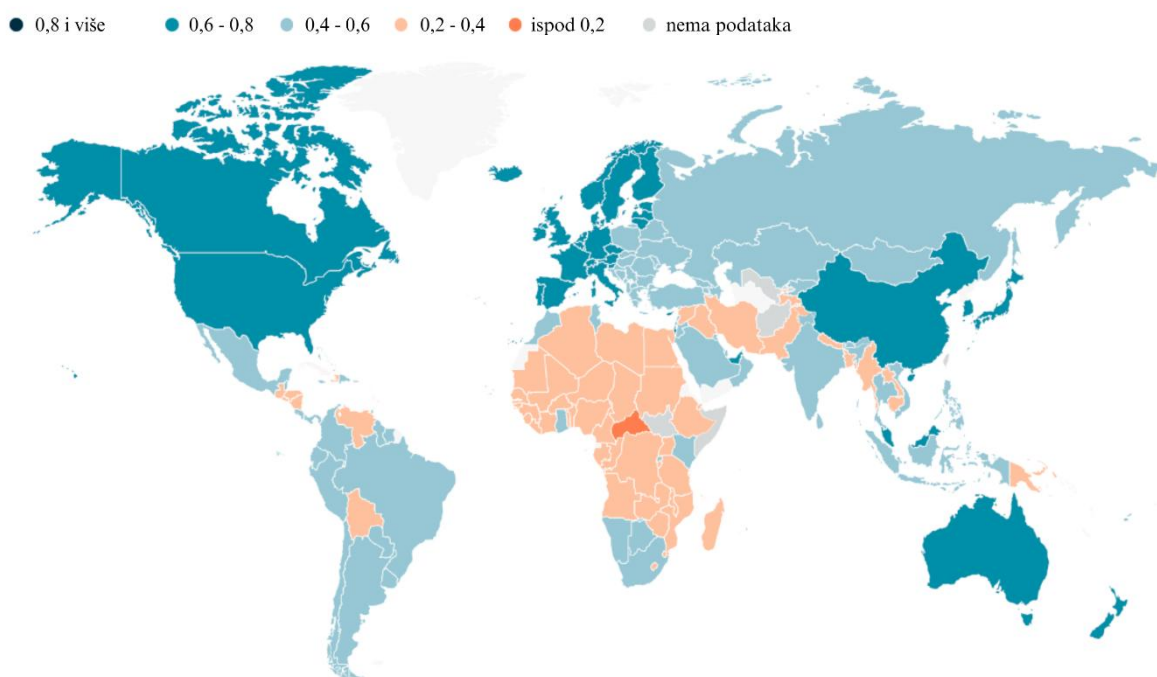
Osim spomenutoga, u Republici Hrvatskoj ne postoji opsežnije istraživanje o prihvaćanju i korištenju generativne umjetne inteligencije, a posebno bi zanimljivo bilo usporediti ispitanike iz Republike Hrvatske sa rezultatima ispitanika iz Ujedinjenog Kraljevstva, europske zemlje koja ne podliježe direktno regulativama EU akta o umjetnoj inteligenciji, ali i zemlje koja doživljava ogroman industrijski rast u razvoju umjetne inteligencije te će uskoro postati treće najveće središte na svijetu za umjetnu inteligenciju, nakon SAD-a i Kine (Browne, 2024). Službeni dokumenti vlade Ujedinjenog Kraljevstva navode kako oni već jesu treće najveće tržište u svijetu kada je umjetna inteligencija u pitanju (Vlada UK, 2025).

Međunarodni monetarni fond (IMF – *International Monetary Fund*) jedan je od globalno priznatih i relevantnih izvora podataka za razna stručna i akademska istraživanja. U okviru svojih baza podataka IMF navodi i Indeks spremnosti na umjetnu inteligenciju (AIPI) koji procjenjuje razinu spremnosti na umjetnu inteligenciju u 174 zemlje, na temelju bogatog skupa makrostrukturnih pokazatelja koji pokrivaju digitalnu infrastrukturu zemalja, ljudski kapital i politike tržišta rada, inovacije i ekonomsku integraciju te regulativu i etiku. Indeks o spremnosti na umjetnu inteligenciju uključuje službene podatke, istraživanja pouzdanih podataka i istraživanja percepcija koje je prikupilo osam institucija: Fraser institut, Međunarodna organizacija rada, Međunarodna unija za telekomunikacije, Ujedinjeni narodi, Konferencija Ujedinjenih naroda za trgovinu i razvoj, Svjetski poštanski savez i Svjetska banka i Svjetski gospodarski forum. AIPI je zapravo zbroj četiri ključne dimenzije: digitalna infrastruktura, ljudski kapital, tehnološke inovacije i pravni okviri (IMF, 2024).

IMF (2024) je 174 zemlje podijelio u tri glavne kategorije:

- Napredne ekonomije (ADVEC – engl. *Advanced economies*)
- Razvijajuće ekonomije (EME – engl. *Emerging market economies*)
- Zemlje s niskim dohotkom (LIC – engl. *Low-income countries*).

Prema podacima koje navodi IMF (2024), zemlja koja je najspremnija na umjetnu inteligenciju jest Singapur (0,80), a slijede ju Danska (0,78), SAD i Nizozemska (0,77), Estonija, Finska i Švicarska (0,76), Novi Zeland, Njemačka i Švedska (0,75), Luksemburg (0,74), Japan, UK, Australija, Južna Koreja i Izrael (0,73), Austrija (0,72), Kanada i Noveška (0,71), Hong Kong, Island i Francuska (0,70). Na začelju ovog popisa nalaze se Južni Sudan (0,11), Afganistan (0,13), Republika Centralna Afrika (0,18) i Somalija (0,2). Republika Hrvatska s vrijednošću od 0,58 na Indeksu spremnosti na umjetnu inteligenciju nalazi se u društvu Grčke, Rumunjske, Bugarske i Saudijske Arabije koje dijele istu vrijednost. Premda Republiku Hrvatsku dijeli samo 0,03 postotna poena od prve kategorije (napredne ekonomije), ona nije u društvu najrazvijenijih svjetskih ekonomija već je po vrijednosti na indeksu u društvu zemalja istočne Europe.



Slika 15. Karta svijeta po IMF-ovom Indeksu spremnosti na umjetnu inteligenciju

Izvor: IMF (2024)

Ukoliko se IMF-ov Indeks spremnosti na umjetnu inteligenciju razdvoji na dimenzije koje ga čine, kvantificirane vrijednosti po pojedinoj dimenziji za promatranu Republiku Hrvatsku i Ujedinjeno Kraljevstvo su sljedeće:

Tablica 14. Usporedba Republike Hrvatske i Ujedinjenog Kraljevstva po dimenzijama
Indeksa spremnosti na umjetnu inteligenciju

Zemlja	Digitalna infrastruktura	Inovacije i ekonomske integracije	Ljudski kapital i politike tržišta rada	Regulacije i etika	AIPI
Hrvatska	0,17	0,14	0,13	0,14	0,58
UK	0,18	0,16	0,17	0,21	0,73

Izvor: IMF (2024)

Kako je vidljivo iz tablice 14., Ujedinjeno Kraljevstvo naprednije je od Republike Hrvatske u svim promatranim dimenzijama, ali razlike su manje kod digitalne infrastrukture te inovacija i ekonomskih integracija, a nešto su ipak izraženije kod ljudskog kapitala i politika tržišta rada te regulacije i etike po pitanju umjetne inteligencije (IMF, 2024).

Nadalje, Oxford Insights (2023) kreirao je Indeks spremnosti državnih vlada na umjetnu inteligenciju u kojem su usporedili strategije i aktivnosti koje vlade diljem svijeta rade kako bi umjetnu inteligenciju približile svojim građanima i implementirale ju u pružanje javnih usluga. Navedeni indeks skrojen je od 39 indikatora unutar deset dimenzija, od čega su tri ključne dimenzije – vladine mjere, tehnološki sektor i digitalna infrastruktura.

Na temelju navedenih parametara, Oxford Insights (2023) navodi kako su Sjedinjene Američke Države najspremnije (84,8), a slijede ih Singapur (81,97), UK (78,57), Finska (77,37), Kanada (77,07), Francuska (76,07), Južna Koreja (75,65) i Njemačka (75,26). Hrvatska se u tom pogledu nalazi na 70. mjestu u svijetu (49,34), što je lošiji rezultat u odnosu na izvješća Oxford Insightsa iz 2022. godine kada je Hrvatska bila na 66. mjestu, 2021. godine kada je Hrvatska bila na 61. mjestu ili 2020. godine kada je Hrvatska bila na 58. mjestu. Američka tehnološka tvrtka Cisco također ima svoj Indeks spremnosti na umjetnu inteligenciju. Ciscov Indeks spremnosti na umjetnu inteligenciju obuhvatio je 30 tržišta, među kojima nije Republika Hrvatska. Ipak, za UK ističu kako je zemlja po mnogim kriterijima spremna na integraciju umjetne inteligencije, no ističu im nedostatke slabog privlačenja talenata (Cisco, 2023).

Za razliku od Republike Hrvatske koja umjetnu inteligenciju vidi kao samo jedan od mnoštva elemenata koji će pomoći digitalnom razvoju države (Središnji državni ured za razvoj digitalnog društva, 2023), Ujedinjeno Kraljevstvo ima Nacionalnu strategiju za umjetnu inteligenciju koju su objavili još 2021. godine (*Secretary of State for Digital, Culture, Media and Sport of UK*, 2021). Prva rečenica u navedenom planu glasi – „Naš desetogodišnji plan je

da učinimo Britaniju globalnom AI supersilom“. Vlada Ujedinjenog Kraljevstva prepoznala je umjetnu inteligenciju kao iznimno važnu granu njihovog gospodarstva, a dokaz tomu je i potvrda ministra tehnologije UK-a koji navodi kako britanske tvrtke koje se bave umjetnom inteligencijom privlače u prosjeku 200 milijuna funti dnevno. Osim toga, Vladin plan predviđa akcijski plan za razvoj mogućnosti umjetne inteligencije od 14 milijardi funti te val ulaganja od 25 milijardi funti u infrastrukturu podatkovnih centara u UK-u (*Department for Science, Innovation and Technology*, 2025). Microsoftovo izvješće ukazuje na potencijal vrijedan 550 milijardi britanskih funti kojim bi umjetna inteligencija povećala BDP Ujedinjenog Kraljevstva u narednih 10 godina (Dawson, 2024).

Za razliku od Ujedinjenog Kraljevstva, koje više nije zemlja članica Europske unije, Hrvatska se snažno oslanja na programe koje EU potiče i sufinancira, kao što su Next Generation EU te Instrument za oporavak i otpornost koji su odvojili 4,4 milijarde eura za projekte povezane s umjetnom inteligencijom. Od zemalja Europske unije, Italija i Španjolska su predvodnici po pitanju ulaganja u projekte iz područja umjetne inteligencije, dok su ulaganja Republike Hrvatske u skladu s prosjekom EU (Europski parlament, 2024b). U travnju 2025. godine Europska komisija (2025) objavila je ambiciozan akcijski plan za Europu kao kontinent umjetne inteligencije uz investicijsku inicijativu vrijednu 200 milijardi eura samo za ulaganja u umjetnu inteligenciju (Europska komisija, 2025).

Podatci Nacionalnog instituta za ekonomska istraživanja navode da alate generativne umjetne inteligencije u većoj ili manjoj mjeri koristi 40 % odraslih u SAD-u (Bick i sur., 2024). Podatci globalnog istraživanja koje je Microsoft (2024) proveo u ljeto 2023. godine pokazuju kako 20 % ispitanika aktivno koristi alate generativne umjetne inteligencije, a 18 % ispitanika su eksperimentalni korisnici. Podatci za Republiku Hrvatsku (Kopal i sur., 2024) i Ujedinjeno Kraljevstvo (Deloitte, 2024) su slični, ali ipak za nijansu veći jer su istraživanja provedena nešto kasnije. Upotreba alata generativne umjetne inteligencije detaljnije je pojašnjena u poglavljima 6,1. i 6,2.

5.1. Upotreba generativne umjetne inteligencije u Hrvatskoj

Istraživanje koje je provela Prizma (2024) u suradnji s Hrvatskom poštom, Kufner grupom, Finom i Effectusom, na uzorku od 2217 ljudi, a koje je objavljeno u rujnu 2024. godine, pokazalo je kako u Hrvatskoj raste pismenost po pitanju umjetne inteligencije. Samo 5 % ispitanika izjavilo je kako nije upoznato s umjetnom inteligencijom, 55 % ispitanika čulo je za umjetnu inteligenciju, a 16 % ju je isprobalo i nije više koristilo. Ipak, 20 % ispitanika izjavilo

je kako umjetnu inteligenciju koristi povremeno, 4 % ispitanika izjavilo je kako ju koristi redovito i samo 1 % ispitanika izjavio je da umjetnu inteligenciju koristi svakodnevno. Premda djeluje kao da te brojke nisu velike, one su značajno veće nego u istraživanju iz 2023. godine (Kopal i sur., 2024).

U istom istraživanju samo se 5 % ispitanika izjasnilo kako već plaća neki alata umjetne inteligencije, a njih 35 % se uskoro planira pretplatiti na neki od alata. Rezultati pokazuju i kako 59 % ispitanika opće populacije ima negativne asocijacije na umjetnu inteligenciju, a 31 % ispitanika ima pozitivne asocijacije. (Prizma, 2024; Kopal i sur., 2024).

Kada je u pitanju upotreba alata umjetne inteligencije, 45 % ispitanika nije upoznato s njima, 20 % je vidjelo kako se koristi, ali nije samo koristilo. Dakle, oko 35 % stanovništva RH barem je probalo koristiti neki od alata umjetne inteligencije. Alate umjetne inteligencije svakodnevno koristi tek 2 % ispitanika, 10 % koristi ih na tjednoj bazi, 8 % na mjesečnoj bazi, a 15 % ispitanika tek je probalo koristiti. Od ispitanika koji su barem probali koristiti, polovica njih koristila je alate umjetne inteligencije isključivo u privatne svrhe, 20 % njih isključivo u poslovne svrhe, a preostalih 30 % koristi alate umjetne inteligencije i u privatnom i u poslovnom životu (Prizma, 2024; Kopal i sur., 2024).

Još neki od zanimljivih zaključaka koji se navodi u istraživanju koje je objavila Prizma (2024):

- Više građana smatra da prilike i prednosti proizvoda i usluga koji koriste umjetnu inteligenciju nadmašuju prijetnje i rizike.
- Više je onih koji ne vjeruju sustavima umjetne inteligencije, bez obzira što smatraju kako će sustavi s vremenom biti sve bolji.
- Umjetna inteligencija većini je prihvatljiva za osnovnu razinu korisničke podrške, ali ne vjeruju da će umjetna inteligencija u potpunosti zamijeniti čovjeka.
- Kao najveći rizik ispitanici percipiraju manipulaciju i zlonamjerno korištenje.
- Gotovo polovica ispitanika smatra da će u idućih pet do deset godina pod utjecajem umjetne inteligencije nestati mnogo poslova u Hrvatskoj.

5.2. Upotreba generativne umjetne inteligencije u Ujedinjenom Kraljevstvu

Kako navodi Deloitte (2024), koji je proveo istraživanje na 4150 odraslih osoba u Ujedinjenom Kraljevstvu, otprilike 36 % stanovništava, ekvivalent 18 milijuna ljudi u dobi od 16 do 75 godina u UK-u barem je probalo koristiti alate generativne umjetne inteligencije. Od onih koji navode da koriste generativnu umjetnu inteligenciju, 10 % njih koristi ju svakodnevno, 26 %

navodi kako ju koristi na tjednoj bazi, a 41 % ispitanika navodi da koristi rjeđe od jednom mjesečno.

Nadalje, Deloitte (2024) navodi kako je 60 % stanovnika UK-a svjesno generativne umjetne inteligencije, ali samo 36 % ju je barem probalo koristiti. Ono što je zanimljivo kod primjene alata generativne umjetne inteligencije jesu demografske karakteristike korisnika gdje se vidi jaz po spolu i dobi. Istraživanje pokazuje kako 43 % muškaraca koristi generativnu umjetnu inteligenciju, u usporedbi sa samo 28 % žena. Navedenu tehnologiju primarno koriste mlađi korisnici, čak 62 % njih u dobi od 16 do 34 godine, a samo 14 % osoba u dobi od 55 do 75 godina.

Deloitte (2024) navodi kako samo 14 % stanovnika UK-a koristi alate generativne umjetne inteligencije za posao, ali od njih čak 74 % navodi kako im alati generativne umjetne inteligencije značajno povećavaju produktivnost. Nadalje, ispitanici ističu kako navedenu tehnologiju najviše koriste za generiranje ideja, traženje informacija, stvaranje tekstualnih sadržaja i pisanje, pisanje e-pošte i sažimanje tekstova.

KMPG (2025) navodi kako 69 % ljudi u UK-u koristi neku vrstu umjetne inteligencije za posao, studiranje ili osobne potrebe, a tek 42 % ljudi navodi kako su spremni vjerovati umjetnoj inteligenciji. Većina britanskih ispitanika zabrinuta je oko negativnih posljedica umjetne inteligencije i smatra kako je potrebna snažnija regulacija navedene tehnologije.

6. KONCEPTUALNI DIZAJN ISTRAŽIVANJA

Venkatesh i sur. (2012) navode kako su sve stavke u istraživačkom instrumentu izvornog UTAUT2 rada mjerene koristeći sedmostupanjsku Likertovu ljestvicu, s krajnjim vrijednostima „izrazito se ne slažem“ (engl. *Strongly disagree*) i „izrazito se slažem“ (engl. *Strongly agree*). Dob je u radu mjerena u godinama, a koncept korištenja tehnologije mjeran je kao formativni kompozitni indeks raznolikosti i učestalosti korištenja tehnologije, u slučaju rada Venkatesh i sur. (2012) to je bila upotreba mobilnog interneta u različite svrhe. Navedeni autori ponudili su popis šest popularnih internetskih aplikacija i od ispitanika se tražilo da navedu učestalost upotrebe za svaku od tih aplikacija pri čemu su ponovno koristili sedmostupanjsku Likertovu ljestvicu s krajnjim vrijednostima od „nikada“ do „više puta u danu“. Venkatesh i sur. (2012) navode kako su po uzoru na Sharma i sur. (2009) za mjerenje korištenja tehnologije koristili ljestvice potkrijepljene ponašanjem koje mogu biti podložne relativno visokoj varijanci zajedničke metode (CMV).

Mjerne ljestvice unutar ovog istraživanja, kako je i prethodno spomenuto, sastoje se od sedam konstrukta koje navode Venkatesh i sur. (2012) u izvornom radu o UTAUT2 modelu te je model proširen s konstruktom povjerenja (Gansser & Reich, 2021; Foroughi i sur., 2023) i konstruktom osobne inovativnosti (Mohd Rahim i sur., 2022), a izmijenjene su i čestice konstrukta USE (korištenje) po uzoru na Venkatesh i sur. (2008). Osim čestica navedenih konstrukta model sadrži i osam demografskih pitanja, pri čemu su tri usmjerena na moderatore iskustva, spola i dobi, a ostala demografska pitanja stavljena su radi boljeg razumijevanja ciljane skupine ispitanika.

U okviru ovog doktorskog rada rada korištena je Likertova ljestvica od 1 do 7, pri čemu su krajnje vrijednosti 1 – „U potpunosti se ne slažem“, a 7 – „U potpunosti se slažem“. Spol je kodiran kao 1 ili 2 varijabla, pri čemu 1 predstavlja muški spol, a 2 ženski spol. Iskustvo je ispitivano mjernim razredima o mjesecima korištenja alata generativne umjetne inteligencije, pri čemu je (1) kraće od mjesec dana, (2) 1 – 3 mjeseca, (3) 3 – 6 mjeseci, (4) 6 – 12 mjeseci, (5) duže od 12 mjeseci. Također, u istraživanje je stavljeno i pitanje „Koristite li plaćenu verziju kod nekog od navedenih alata?“, pri čemu su ispitanici mogli odgovoriti sa „Da“ ili „Ne“. Navedena demografska karakteristika prethodno nije korištena u radovima kao moderator te zbog toga nije stavljena kao moderator niti u ovom radu, ali će se detaljno prokomentirati u raspravi.

Osim spomenute Likertove ljestvice i čestica konstrukta, radi usporedivosti, po uzoru na istraživanje Venkatesh i sur. (2012) preuzeta su i mjerenja za moderatore pa su tako ponuđene opcije kod spola muškarci i žene, dob je mjerena u godinama, a iskustvo u mjesecima korištenja navedene tehnologije u središtu istraživanja.

Ono što je jedna od ključnih razlika ovoga istraživanja u metodološkom smislu, u odnosu na izvorni rad (Venkatesh i sur., 2012), jest u činjenici da je izvorni rad provodio prikupljanje u dvije etape, a unutar ovoga istraživanja prikupljanje podataka bilo je u samoj jednoj etapi.

6.1. Instrument istraživanja

Instrument istraživanja korišten u ovom doktorskom radu visoko je strukturirana anketa, izrađena u Alchemeru, platformi za izradu i provedbu istraživanja. Navedeni softver koristi se za provođenje kvantitativnog istraživanja CAWI metodom (engl. *Computer-assisted web interviewing*). Anketa je izrađena na hrvatskom i engleskom jeziku s ciljem provedbe istraživanja u Republici Hrvatskoj i Ujedinjenom Kraljevstvu.

Anketa se sastoji od ukupno 46 pitanja, pri čemu je osam demografskih pitanja i 38 pitanja koja ispituju odrednice prihvaćanja i korištenja generativne umjetne inteligencije.

Prije provođenja glavnog istraživanja, provedeno je pilot-istraživanje na 23 ispitanika, dominantno među poznatim kontaktima koji se uklapaju u ciljnu skupinu, a imaju potrebno predznanje o tematici, s ciljem ispitivanja razumijevanja pitanja i konteksta varijabli. Također, kroz platformu Alchemer korištene su različite funkcionalnosti osiguravanja kvalitete podataka poput nasumičnog redoslijeda prikazivanja ponuđenih odgovora, formatiranja traženih odgovora, testiranja upotrebljivosti i prikaza iz pozicije korisnika te drugih srodnih funkcionalnosti.

Tablica 15. Istraživački instrument na hrvatskom i engleskom jeziku

HRVATSKI	ENGLISH
Spol	Gender
Muško Žensko	Male Female
Godina Vašeg rođenja	Year of birth
Upišite: _____	Write: _____
Označite alate generativne umjetne inteligencije koju ste koristili:	Mark the tools of Gen AI that you have at least tried to use:

Adobe Firefly AlphaCode ChatGPT Claude Copilot Dall-E DeepSeek Elicit Gemini GitHub Copilot Google NotebookLM Grammarly Grok HeyGen Ideogram Jasper AI Leonardo Llama Midjourney Mistral Perplexity Poe Quillbot Runway SciteAI Stability AI Stable Diffusion Ostalo – upišite: _____	Adobe Firefly AlphaCode ChatGPT Claude Copilot Dall-E DeepSeek Elicit Gemini GitHub Copilot Google NotebookLM Grammarly Grok HeyGen Ideogram Jasper AI Leonardo Llama Midjourney Mistral Perplexity Poe Quillbot Runway SciteAI Stability AI Stable Diffusion Other – Write In _____
Koliko dugo koristite alate generativne umjetne inteligencije?	Have long have you been using generative AI tools?
Manje od 1 mjeseca. 1 – 3 mjeseca. 3 – 6 mjeseci. 6 – 12 mjeseci. Duže od 12 mjeseci.	Less than 1 month. 1 – 3 months. 3 – 6 months. 6 – 12 months. More than 12 months.
Koristite li plaćenu verziju kod nekog od navedenih alata?	Are you using the paid version of any of the above tools?
Da Ne	Yes No
Jeste li u radnom odnosu?	Do you work?
Idem u srednju školu. Studiram. Studiram i radim. Radim. Niti radim, niti studiram.	I am a high-school student. I am a university / college student. I am student but I work. Yes, I do work. I do not work, and I am not student.
Mjesto stanovanje?	I live in

Ruralno područje Mali grad (do 15 000 stanovnika) Srednji grad (do 50 000 stanovnika) Veliki grad (preko 50 000 stanovnika)	Rural area (less than 10,000 people) Small town (10,000 – 75,000 people) Large town or small city (75,000 – 250,000 people) Major city (more than 250,000 people)
Najviši završeni stupanj obrazovanja	Highest education level completed
Osnovna škola ili niže Srednja škola Viša / Stručni studij Prije diplomski studij Diplomski studij Poslijediplomski studij	Primary school or below Secondary school (e.g. GED / GCSE) High school diploma / A-levels Technical / community college Undergraduate degree (BA / BCs / other) Postgraduate (MA / MSc / Mphil / other) Doctorate degree (PhD / other)
Očekivani učinak	Performance Expectancy
• Smatram da je korištenje alata generativne umjetne inteligencije korisno. • Korištenje alata generativne umjetne inteligencije povećava moje šanse da ostvarim stvari koje su mi bitne. • Korištenje alata generativne umjetne inteligencije pomaže mi da brže obavim aktivnosti. • Korištenje alata generativne umjetne inteligencije povećava moju produktivnost.	• I find Gen AI useful in my daily life. • Using Gen AI increases my chances of achieving things that are important to me. • Using Gen AI helps me accomplish things more quickly. • Using Gen AI increases my productivity.
Očekivani napori	Effort Expectancy
• Lako mi je naučiti koristiti alate generativne umjetne inteligencije. • Moja interakcija s alatima generativne umjetne inteligencije je jasna i razumljiva. • Smatram da su alati generativne umjetne inteligencije jednostavni za korištenje. • Lako mi je postati vješt u korištenju generativne umjetne inteligencije.	• Learning how to use Gen AI is easy for me. • My interaction with Gen AI is clear and understandable. • I find Gen AI easy to use. • It is easy for me to become skillful at using Gen AI.
Društveni utjecaj	Social Influence
• Osobe koje su mi važne misle da bih trebao koristiti alate generativne umjetne inteligencije. • Osobe koje utječu na moje ponašanje misle da bih trebao koristiti alate generativne umjetne inteligencije.	• People who are important to me think that I should use Gen AI. • People who influence my behavior think that I should use Gen AI. • People whose opinions that I value prefer that I use Gen AI.

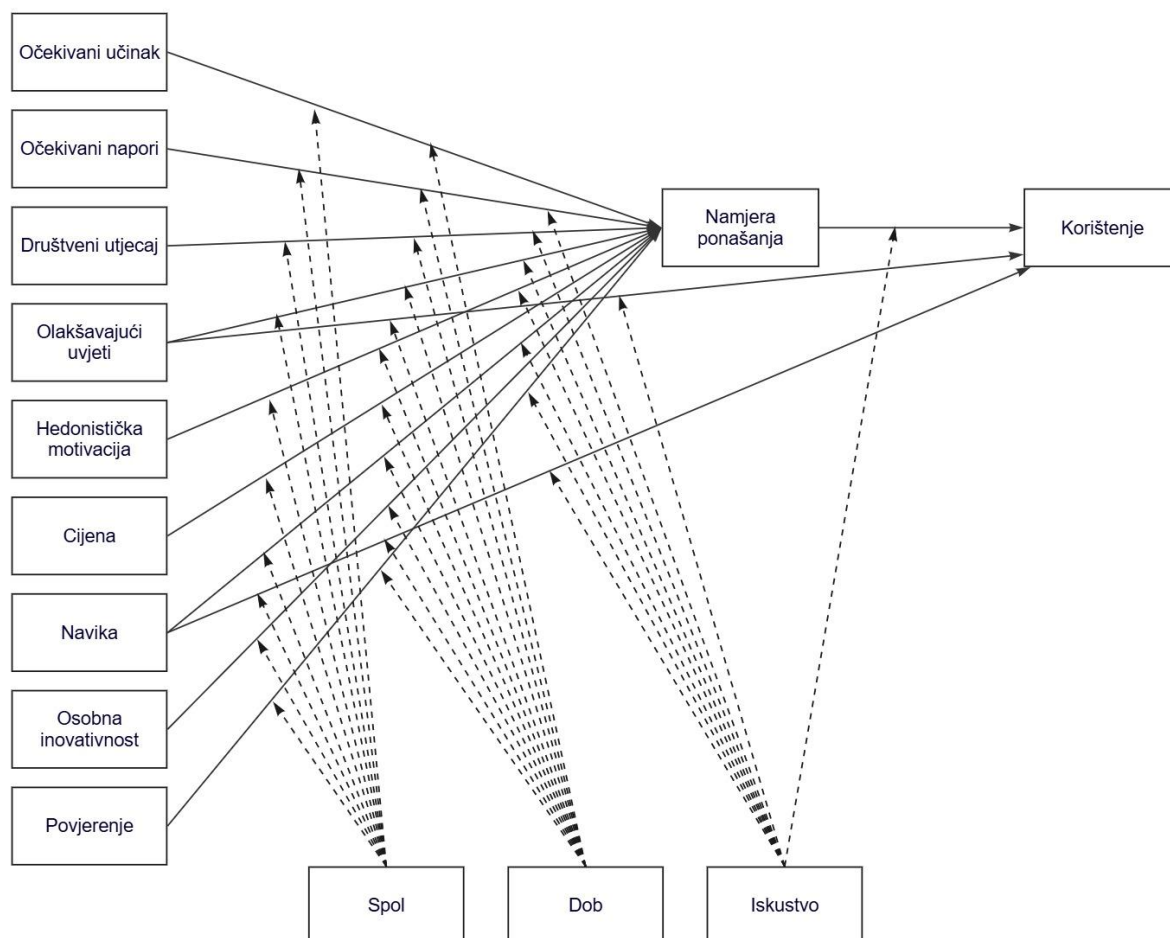
Osobe čije mišljenje cijenim preferiraju da koristim alate generativne umjetne inteligencije.	
Olakšavajući uvjeti	Facilitating conditions
- Imam potrebne resurse za korištenje alata generativne umjetne inteligencije. - Imam potrebno znanje za korištenje generativne umjetne inteligencije. - Alati generativne umjetne inteligencije kompatibilni su s drugim tehnologijama koje koristim. - Mogu dobiti pomoć od drugih kada imam poteškoća s korištenjem alata generativne umjetne inteligencije.	- I have the resources necessary to use Gen AI. - I have the knowledge necessary to use Gen AI. - Gen AI tools are compatible with the technology I use. - I can get help from other when I have difficulties using Gen AI.
Hedonistička motivacija	Hedonic motivation
- Korištenje alata generativne umjetne inteligencije je zabavno. - Korištenje alata generativne umjetne inteligencije je ugodno. - Korištenje alata generativne umjetne inteligencije je vrlo zanimljivo.	- Using Gen AI is fun. - Using Gen AI is enjoyable. - Using Gen AI is very entertaining.
Cijena	Price value
- Alati generativne umjetne inteligencije pružaju dobru vrijednost za novac. - Alati generativne umjetne inteligencije imaju razumnu cijenu. - Po trenutnoj cijeni, alati generativne umjetne inteligencije pružaju dobru vrijednost.	- Gen AI tools are reasonably priced. - Gen AI is a good value for money. - At the current prices, Gen AI tools provide a good value.
Navika	Habit
- Korištenje alata generativne umjetne inteligencije postalo mi je navika. - Ovisan sam o korištenju alata generativne umjetne inteligencije. - Moram koristiti alate generativne umjetne inteligencije. - Korištenje alata generativne umjetne inteligencije postalo mi je prirodno.	- The use of Gen AI has become a habit for me. - I am addicted to using Gen AI. - I must use Gen AI. - Using Gen AI has become natural to me.
Osobna inovativnost	Personal innovativeness
- Kada čujem za novu tehnologiju potražim način da je odmah isprobam. - Među svojim prijateljima, ja sam uglavnom prvi/a koji proba nove tehnologije. - Volim eksperimentirati s novim tehnologijama.	- If I hear about new technology, I would look for ways to experiment with it. - Among my peers, I am usually the first to try out new technologies. - I like to experiment with new technologies.

Povjerenje	Trust
• Koristit ću alate generativne umjetne inteligencije ako osjetim da je sadržaj pouzdan. • Koristit ću alate generativne umjetne inteligencije ako osjetim da mi pruža pouzdane informacije. • Koristit ću alate generativne umjetne inteligencije ako osjetim da ispunjava moja očekivanja. • Koristit ću alate generativne umjetne inteligencije ako osjetim da je sigurna.	• I will use Gen AI if I feel that the content is trustworthy. • I will use Gen AI if I feel that Gen AI provides reliable information. • I will use Gen AI if I feel that Gen AI meets my expectations. • I will use Gen AI if I feel that Gen AI is secure.
Namjera ponašanja	Behavioral intention
• Namjeravam nastaviti koristiti alate generativne umjetne inteligencije u svakodnevnom životu. • Uvijek ću nastojati koristiti alate generativne umjetne inteligencije u svakodnevnom životu. • Planiram često koristiti alate generativne umjetne inteligencije.	• I intend to continue using Gen AI in the future. • I will always try to use Gen AI in my daily life. • I plan to continue using Gen AI frequently.
Korištenje	USE
Trajanje: „U prosjeku, koliko sati tjedno koristite alate generativne umjetne inteligencije?“ Ljestvica: 1 = Manje od 1 sat 2 = 1–2 sata 3 = 2–3 sata 4 = 3–4 sata 5 = 4–5 sati 6 = 5–6 sati 7 = Više od 6 sati Učestalost: „Koliko često koristite alate generativne umjetne inteligencije?“ Ljestvica: 1 = Nekoliko puta godišnje ili rjeđe 2 = Jednom mjesečno 3 = Nekoliko puta mjesečno 4 = Jednom tjedno 5 = Nekoliko puta tjedno 6 = Jednom dnevno 7 = Nekoliko puta dnevno	Duration: On average, how many hours per week do you use the Gen AI tools?" Scale: 1 = Less than 1 hour 2 = 1–2 hours 3 = 2–3 hours 4 = 3–4 hours 5 = 4–5 hours 6 = 5–6 hours 7 = More than 6 hours Frequency: How often do you use Gen AI tools? Scale: 1 = A few times a year or less 2 = Once a month 3 = A few times a month 4 = Once a week 5 = A few times a week 6 = Once a day 7 = A few times a day or more

Intenzitet: „Kako procjenjujete intenzitet Vaše trenutačne uporabe alata generativne umjetne inteligencije?“ Ljestvica: 1 = Izuzetno niska uporaba 2 = Vrlo niska uporaba 3 = Niska uporaba 4 = Umjerena uporaba 5 = Visoka uporaba 6 = Vrlo visoka uporaba 7 = Intenzivna uporaba	Intensity: How do you consider the extent of your current use of Gen AI tools? Scale: 1 = Extremely low use 2 = Very low use 3 = Low use 4 = Moderate use 5 = High use 6 = Very high use 7 = Heavy use
--	--

6.2. Prošireni UTAUT2 model

Kako je ranije rečeno, za potrebe ovoga rada, UTAUT2 je proširen i prilagođen za istraživanje u kontekstu ispitivanja prihvatanja i korištenja generativne umjetne inteligencije, a prikaz svih odnosa i hipoteza prikazan je na slici 16., ispod teksta.



Slika 16. Prikaz konceptualnog modela – prošireni UTAUT2 model

Izvor: izrada autora

Kako je prikazano na slici 16., u modelu se ispituje 41 utjecaj, uključujući sve odnose između zavisnih i nezavisnih varijabli i njihovih moderatora. Na temelju tih odnosa napravljene su i hipoteze te istraživačko pitanje.

6.3. Hipoteze

Rad sadržava ukupno 12 hipoteza, odnosno ispituje se 10 odnosa između zavisnih i nezavisnih varijabli te moderatora, ali su hipoteza 4 i hipoteza 7 podijeljene na dvije podhipoteze jer model predviđa kako zavisne varijable olakšavajući uvjeti (FC) i navika (HT) utječu i na namjeru ponašanja (BI), ali i na stvarno ponašanje (USE). S obzirom na to da pojedine hipoteze mjere i po četiri odnosa, potencijalno očekivani scenarij je potvrđivanje tek dijela tih odnosa. U takvim situacijama prvenstveno se promatra odnos između zavisne i nezavisne varijable, ukoliko je on potvrđen, a neki od moderatora nije – hipoteza se djelomično prihvaća. Ipak, ukoliko se moderatori pokažu značajnim, a odnos između zavisne i nezavisne varijable ne bude – hipoteza se ne prihvaća. U konačnici, napravljeno je i istraživačko pitanje koje objedinjuje cjelokupni konceptualni model istraživanja.

***H1:** Dob i spol moderiraju utjecaj očekivanog učinka na namjeru ponašanja, tako da je utjecaj jači kod mlađih muškaraca.*

***H2:** Dob, spol i iskustvo moderiraju utjecaj očekivanih napora na namjeru ponašanja, tako da je utjecaj jači kod mlađih žena u ranoj fazi iskustva s tehnologijom.*

***H3:** Dob, spol i iskustvo moderiraju utjecaj društvenog utjecaja na namjeru ponašanja, tako da je utjecaj jači kod starijih žena u ranoj fazi iskustva s tehnologijom.*

***H4a:** Dob, spol i iskustvo moderiraju utjecaj olakšavajućih uvjeta na namjeru ponašanja, tako da je utjecaj jači kod starijih žena u ranoj fazi iskustva s tehnologijom.*

***H4b:** Dob i iskustvo moderiraju utjecaj olakšavajućih uvjeta na korištenje, tako da je utjecaj jači kod starijih, posebno onih s većom razinom iskustva s tehnologijom.*

***H5:** Dob, spol i iskustvo moderiraju utjecaj hedonističke motivacije na namjeru ponašanja, tako da je utjecaj jači među mlađim muškarcima u ranoj fazi iskustva s tehnologijom.*

H6: *Dob i spol moderiraju utjecaj vrijednosti cijene na namjeru ponašanja, tako da je utjecaj jači među starijim ženama.*

H7a: *Dob, spol i iskustvo moderiraju utjecaj navike na namjeru ponašanja, tako da je utjecaj jači kod starijih muškaraca s većom razinom iskustva s tehnologijom.*

H7b: *Dob, spol i iskustvo moderiraju utjecaj navike na korištenje tehnologije, tako da je utjecaj jači kod starijih muškaraca s većom razinom iskustva s tehnologijom.*

H8: *Dob i spol moderiraju utjecaj osobne inovativnosti na namjeru ponašanja, tako da je utjecaj jači kod mlađih muškaraca.*

H9: *Dob, spol i iskustvo moderiraju utjecaj povjerenja na namjeru ponašanja, tako da je utjecaj jači među mlađim muškarcima s većom razinom iskustva s tehnologijom.*

H10: *Iskustvo moderira utjecaj namjere ponašanja na korištenje tehnologije, tako da je utjecaj jači kod korisnika s manje iskustva.*

Istraživačko pitanje: *Je li i u kojoj mjeri moguće predvidjeti namjeru ponašanja i stvarno ponašanje korisnika generativne umjetne inteligencije na temelju predloženog konceptualnog modela koji je utemeljen na UTAUT2 modelu (očekivani učinak, očekivani naponi, društveni utjecaj, olakšavajući uvjeti, hedonistička motivacija, cijena i navika moderirani dobi, spolom i iskustvom) te proširen konstruktima osobne inovativnosti i povjerenjem?*

S ciljem provjere postavljenih hipoteza i istraživačkog pitanja, provedeno je i empirijsko istraživanje čiji će rezultati biti detaljno analizirani i prikazani u nastavku rada, pružajući sveobuhvatan pregled empirijskih dokaza koji podupiru ili opovrgavaju postavljene pretpostavke te daju odgovor na istraživačko pitanje.

7. EMPIRIJSKO ISTRAŽIVANJE

Kako sugerira ranije navedeni teorijski dio ovog doktorskog rada, razumijevanje odrednica koje utječu na prihvatanje i korištenje generativne umjetne inteligencije vrlo je aktualna tema unutar koje postoje nedovoljno istražena područja.

7.1. Metodologija primarnog istraživanja

Venkatesh i sur. (2012) navode kako su za analizu podataka koristili metodu parcijalnih najmanjih kvadrata (PLS) (engl. *Partial least squares*), odnosno PLS-SEM, jer istraživački nacrt sadržava veći broj interakcijskih odnosa, što značajno usmjerava metodološki pristup. Za provjeru pouzdanosti i valjanosti podataka korišten je test unutarnjih konzistentnih pouzdanosti (ICL) (engl. *Internal consistency reliabilities*), a što podrazumijeva Cronbachovu alfu i kompozitnu pouzdanost, ali i test prosječne ekstrahirane varijance (AVE) (engl. *Average variance extracted*) za provjeru konvergentne valjanosti konstrukta. Također, navedeni autori napominju kako se za daljnja testiranja multikolinearnosti izračunavaju faktori inflacije varijance (VIF) (engl. *Variance inflation factors*), a kojima prethodi pregled korelacija među egzogenim konstruktima te centriranje svih varijabli korištenih za stvaranje interakcijskih odnosa kakav je korišten i kod izvornog UTAUT-a (Venkatesh i sur. 2003). Modeliranje strukturalnim jednadžbama temeljeno na kovarijanci (CB-SEM - *Covariance-based Structural Equation Modeling*) metoda je strukturalnog modeliranja prvenstveno namijenjena u konfirmatorne svrhe, odnosno testiranje teorije. S druge strane, PLS-SEM je orijentiran na predviđanje te se preporuča za korištenje u eksplanatornim i konfirmatornim istraživanjima (Hair i sur., 2017; Henseler, 2018; Vuković, 2022). Najznačajnija prednost PLS-SEM-a je što se s njim mogu specificirati i formativni i reflektivni mjerni modeli, dok je CB-SEM ograničen samo na reflektivne modele (Dash & Paul, 2021). Kada je cilj istraživanja predviđanje i objašnjavanje ciljnih konstrukta ili identificiranje ključnih prediktora konstrukta preporučuje se koristiti PLS-SEM. Nadalje, PLS-SEM koristan je pri eksplorativnim istraživanjima i kod proširivanja postojeće teorije (Vuković, 2022; Hair i sur., 2012; Hair i sur., 2017a; Mohamed i sur., 2019). Studija koju su proveli Dash i Paul (2021) utvrdila je kako su i CB-SEM i PLS-SEM bili jednako učinkoviti za razvoj strukturalnih odnosa. Premda autori češće koriste CB-SEM nego li PLS-SEM, nakon novijih empirijskih analiza i najnovije literature, može se zaključiti da PLS-SEM pruža više fleksibilnosti za istraživanje i eksperimentiranje s brojnim konfiguracijama. PLS-SEM metoda široko je rasprostranjena te njena upotreba u društvenim

znanostima ubrzano raste zahvaljujući prednostima i mogućnostima koje pruža (Vuković, 2022; Hanafiah, 2020).

Još jedna specifična karakteristika PLS-SEM analize jest činjenica da PLS-SEM ne pretpostavlja kako su podaci normalno distribuirani, što implicira da se parametrijski testovi značajnosti, kakvi su karakteristični za regresijske analize, ne mogu primijeniti za provjeru toga jesu li koeficijenti poput vanjskih opterećenja, vanjske težine i koeficijenti puta značajni. Umjesto toga, kako bi se testirala značajnost procijenjenih koeficijenata putanje PLS-SEM koristi *bootstrapping*, neparametrijsku proceduru (Efron & Tibshirani, 1993; Davison & Hinkley, 1997). *Bootstrapping* stvara poduzorke s nasumično izvučenim opažanjima iz izvornog skupa (sa zamjenom) i zatim se poduzorak koristi za procjenu značajnosti putanja u PLS-SEM modelu. Takav se postupak ponavlja sve dok se ne stvori veliki broj nasumičnih poduzoraka (Davison & Hinkley, 1997; Ringle i sur., 2024). Hair i sur. (2022) preporučuju 5000 *bootstrap* uzoraka. PLS-SEM softverski programi, poput SmartPLS-a, prikazuju i p-vrijednosti (tj. vrijednosti vjerojatnosti), koje predstavljaju vjerojatnost dobivanja empirijske t-vrijednosti barem jednako ekstremne kao one koja je zapravo opažena, pod uvjetom da nul-hipoteza vrijedi. P-vrijednost odgovara na pitanje: ako je težina w_1 zapravo jednaka nuli, koja je vjerojatnost da će slučajno uzorkovanje *bootstrappingom* dati t-vrijednost od najmanje 1,96 (u slučaju razine značajnosti od 5 %, dvosmjernog testa) (Hair i sur., 2022). Kada se izvodi test hipoteze, p-vrijednost pomaže odrediti značajnost rezultata. Hipoteza koja se ispituje naziva se nul-hipoteza, dok je alternativna hipoteza ona koja bi se pretpostavila da je točna ako je nulta hipoteza neistinita. Mala p-vrijednost (obično $\leq 0,05$) ukazuje na jake dokaze protiv nulte hipoteze, odnosno tumači se kako je vrijednost statistički značajna. U literaturi se najčešće koristi razina značajnosti od 5 %, odnosno $p \leq 0,050$. Kao standard testiranja ovog doktorskog rada koristit će takav prag, u skladu s literaturom (Rumsey, 2016; Madjarova i sur., 2022; Greenland i sur., 2016). Nadalje, kod tumačenja rezultata standardiziranih koeficijenata puta (β), Hair i sur. (2022) navode kako se vrijednosti blizu 0 tumače kao izostanak efekta, vrijednosti od 0,10 do 0,30 kao slabi efekt, vrijednosti od 0,30 do 0,50 kao umjereni efekt i vrijednosti preko 0,50 kao snažan efekt. Ipak, kod modela s većim brojem prediktora na zavisnoj varijabli treba uzeti u obzir simultani utjecaj prediktora zbog međusobne povezanosti (kolinearnosti) zbog čega pojedinačni koeficijent puta (β) može biti umanjen jer dijeli objašnjenje s drugim prediktorima.

Kako je prethodno spomenuto, PLS-SEM pogodan je za procjenu reflektivnog mjernog modela, a za čije je konstrukte potrebno napraviti procjenu njihove unutarnje konzistentnosti, ali i konvergentne i diskriminantne valjanosti.

Provjerom pouzdanosti unutarnje konzistentnosti mjeri se stupanj u kojem čestice koje mjere isti konstrukt međusobno koreliraju. Jedna od najčešće korištenih provjera unutarnje konzistentnosti je testiranje Cronbachove alfe (engl. *Cronbach's Alpha*), mjere unutarnje konzistentnosti ljestvice koju je predstavio Lee Cronbach (1951). Riječ je o koeficijentu koji pokazuje u kojoj su mjeri čestice nekog latentnog konstrukta povezane, odnosno mjere li istu latentnu dimenziju. Cronbachova alfa označava se simbolom alfa (α), a računa se na temelju prosječnih korelacija među česticama i varijance ukupnog rezultata, pri čemu vrijednost alfe raste s jačom korelacijom među česticama, ali i s brojem čestica. Što je α bliže 0 to su čestice međusobno nezavisnije, a što je vrijednost bliže 1, to su čestice snažnije povezane (Goforth, 2015). Hair i sur. (2021) navode kako je Cronbachova alfa prilično konzervativna mjera, a to je da se vrijednosti pouzdanosti između 0,60 i 0,70 smatraju prihvatljivima u istraživačkim studijama, vrijednosti između 0,70 i 0,90 kreću se od zadovoljavajućih do dobrih, a vrijednosti iznad 0,95 mogu ukazivati na to da su čestice redundantne, odnosno da su svjesno napuhane korelacije među česticama.

Još jedna vrlo bitna i često korištena mjera za ispitivanje pouzdanosti unutarnje konzistentnosti kod PLS-SEM-a je kompozitna pouzdanost (CR rho_c) (Jöreskog, 1971), za koju vrijede isti pragovi kao i za Cronbachovu alfu. Važno je da vrijednosti budu veće od 0,70 kako bi se smatrale zadovoljavajućim. Kao alternativa konzervativnoj Cronbachovoj alfi i vrlo liberalnoj Jöreskogovoj kompozitnoj pouzdanosti (CR rho_c), Dijkstra (2010) je predstavio kompozitnu pouzdanost (CR rho_a), a koja se obično pozicionira između tih vrijednosti i smatra se prihvatljivim kompromisom (Hair i sur., 2021; Dijkstra, 2014; Dijkstra & Henseler, 2015). Upravo iz navedenih razloga, ove se vrijednosti često izvještavaju zajedno radi potpunije slike, jer Cronbachova alfa može podcijeniti, a CR (rho_c) precijeniti pouzdanost.

Nadalje, kod PLS-SEM analize, konvergentna valjanost reflektivnog mjernog modela procjenjuje se pokazateljem prosječne ekstrahirane varijance – AVE (engl. *Average variance extracted*) ili pokazateljima vanjskih opterećenja (engl. *outer loadings*) svake čestice u konstrukt. AVE pokazatelj definira se kao srednja vrijednost kvadratnih opterećenja čestica povezanih s konstruktom, odnosno zbroj kvadratnih opterećenja podijeljen s brojem indikatora. Upravo iz tog razloga je AVE ekvivalent zajedničkoj varijanci konstrukta. Minimalna prihvatljiva vrijednost AVE pokazatelja je 0,50 (Fornell & Larcker, 1981), što ukazuje na

činjenicu da konstrukt objašnjava 50 % ili više varijance čestica koje ga čine (Hair i sur., 2021). Drugim riječima, to znači da čestice više doprinose konstruktnoj varijanci nego što doprinose pogreškama. S druge strane, vanjska opterećenja ili opterećenja indikatora predstavljaju bivarijantne korelacije između konstrukta i njegovih indikatora. Ona određuju apsolutni doprinos pojedine čestice svom pripadajućem konstrukt i zbog toga su opterećenja od primarne važnosti pri evaluaciji reflektivnih mjernih modela. Ipak, vanjska opterećenja također se interpretiraju i kada su uključene formativne mjere (Hair i sur., 2021). Vrijednost vanjskog opterećenja trebala bi biti iznad 0,70, dok za vrijednosti između 0,4 i 0,7 istraživač treba razmotriti eliminaciju čestice iz mjernog modela (Hair i sur., 2022).

Za provjeru diskriminantne valjanosti na razini konstrukta koristi se Fornell-Larckerov kriterij i HTMT omjer (engl. *Heterotrait-Monotrait ratio*). Fornell i Larcker (1981) razvili su navedeni kriterij za procjenu diskriminantne valjanosti, a što zapravo provjerava jesu li latentne varijable međusobno odvojene i objašnjava li svaki konstrukt jedinstveni dio varijance koji ne dijeli ni s jednim drugim konstruktom u modelu. Fornell i Larcker (1981) predložili su da se izračuna korelacija između svakog para latentnih varijabli te da se usporede sa korijenom pokazatelja AVE svake od njih. Kriterij glasi: korijen pokazatelja AVE svakog konstrukta trebao bi biti veći od bilo koje korelacije tog konstrukta s drugim konstruktom. Premda je Fornell-Larckerov kriterij desetljećima bio standard u provjeri diskriminantne valjanosti, novija metodološka istraživanja pokazala su da ovaj pristup ima manjkavosti (Hair i sur., 2021). Henseler i sur. (2015) proveli su simulacijske studije i utvrdili da Fornell-Larckerov kriterij često ne pronalazi diskriminantne valjanosti. Henseler i sur. (2015) tada su predstavili HTMT omjer koji predstavlja suvremeni pristup provjeri diskriminantne valjanosti, a osmišljen je upravo da nadvlada ograničenja Fornell-Larckerova kriterija. HTMT omjer definira se kao omjer prosječnih korelacija čestica među različitim konstruktima (*heterotrait*) i prosječnih korelacija čestica unutar istog konstrukta (*monotrait*). Drugim riječima, HTMT uspoređuje prosječnu sličnost između čestica koje bi trebale pripadati nekom drugom konstrukt s prosječnom sličnosti čestica koje bi trebale pripadati istom konstrukt. Predloženi prag je da vrijednost HTMT omjera treba biti manja od 0,85 u konzervativnijem pristupu ili 0,90 u liberalnijem pristupu (Henseler i sur., 2015; Voorhees i sur., 2015).

Kako je prethodno spomenuto, za provjeru kolinearnosti među česticama u radu se koristi faktor inflacije varijance – VIF, standardna mjera za procjenu kolinearnosti među indikatorima. Što je VIF vrijednost viša, veća je razina kolinearnosti. Kao najčešći prag u literaturi se uzima vrijednost od 5 i ako je VIF vrijednost veća od 5, to može ukazivati na moguće probleme s

kolinearnosti (Hair i sur., 2021). Kock i Lynn (2012) zagovaraju nešto konzervativniju granicu od 3,3, dok određeni statističari liberalnije gledaju na tu granicu pa toleriraju VIF vrijednosti do 10 (Kutner i sur., 2005; Gujarati, 2003).

7.2. Proces prikupljanja podataka

Kada je riječ o veličini uzorka, PLS-SEM je vrlo učinkovit u radu s malim uzorcima, čak i kada su modeli kompleksni, a to je moguće jer odvojeno računa veze u strukturalnom i u mjernom modelu putem parcijalnih regresija. PLS-SEM u takvim situacijama ima veću snagu, odnosno PLS-SEM metoda ima veću vjerojatnost pronalaska statistički značajnog rezultata ukoliko je on uistinu prisutan u populaciji. Ipak, treba imati na umu da niti jedan statistički test ili metoda ne mogu popraviti uzorak ukoliko je on loš i pretvoriti ga u reprezentativan, prikladan uzorak (Vuković, 2022; Hair i sur., 2011; Hair i sur., 2017a, Hair i sur., 2017b, Hair i sur., 2019).

Hair i sur. (2017b) oponiraju pravilu koje kaže kako bi uzorak trebao biti deset puta veći od broja formativnih indikatora koji se koriste za mjerenje jednog konstrukta ili deset puta veći od najvećeg broja strukturalnih putova usmjerenih prema određenom konstrukt u strukturalnom modelu (Barclay i sur., 1995). Drugim riječima, navedeno pravilo kaže kako bi broj ispitanika trebao biti deset puta veći od maksimalnog broja strelica u modelu. Hair i sur. (2017b) sugeriraju kako se istraživači mogu poslužiti alatima kao što su G*Power kako bi testirali koliko velik uzorak treba biti, ali su predstavili i vlastiti model za odabir veličine uzorka temeljem snage testa (najčešće je to 80 %) koja se želi postići, vidljiv u tablici 16.

Tablica 16. Preporuke veličine uzorka za PLS-SEM pri snazi testa od 80 %

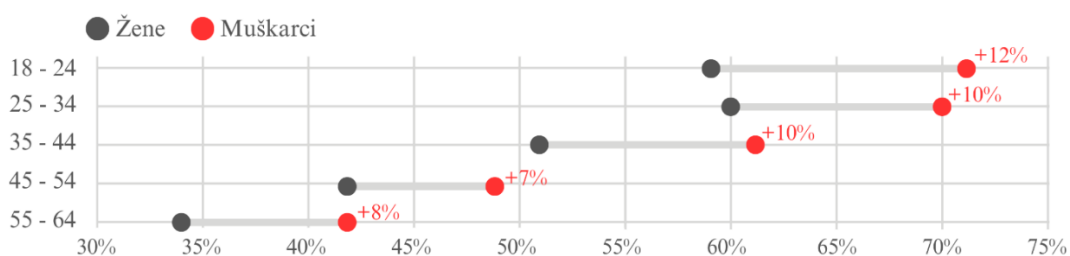
Broj nezavisnih varijabli konstrukta	Razina signifikantnosti											
	10 %				5 %				1 %			
	Minimalni R ²				Minimalni R ²				Minimalni R ²			
	0,10	0,25	0,50	0,75	0,10	0,25	0,50	0,75	0,10	0,25	0,50	0,75
2	72	26	11	7	90	33	14	8	130	47	19	10
3	83	30	13	8	103	37	16	9	145	53	22	12
4	92	34	15	9	113	41	18	11	158	58	24	14
5	99	37	17	10	122	45	20	12	169	62	26	15
6	106	40	18	12	130	48	21	13	179	66	28	16
7	112	42	20	13	137	51	23	14	188	69	30	18
8	118	45	21	14	144	54	24	15	196	73	32	19

9	124	47	22	15	150	56	26	16	204	76	34	20
10	129	49	24	16	156	59	27	18	212	79	35	21

Izvor: izrada autora prema Hair i sur., 2019b; str. 26

Ipak, u ovom primarnom istraživanju, autor se odlučio za prikupljanje kvotnog uzorka na tržištu Ujedinjenog Kraljevstva i Republike Hrvatske. Cilj je bio prikupiti po 400 ispitanika na svakom tržištu, zatim podatke obraditi na cjelovitom uzorku, a potom i obraditi svako tržište zasebno te ih u konačnici i usporediti.

Kvotni uzorak za primarno istraživanje napravljen je po uzoru na istraživanje koje je proveo Oliver Wyman Forum (2024). Navedeno istraživanje provela je jedna od najvećih svjetskih konzultantskih tvrtki, u periodu od lipnja do studenog 2024. godine, na uzorku od gotovo 25 tisuća ispitanika u 16 zemalja, uključujući SAD, Kanadu, Meksiko, Brazil, UK, Francusku, Italiju, Njemačku, Španjolsku, Kinu, Hong Kong, Indiju, Indoneziju, Singapur, UAE i Australiju.



Slika 17. Razlike radno aktivnih korisnika generativne umjetne inteligencije po dobi i spolu

Izvor: prilagođeno prema Svjetski ekonomski forum (2024) prema Oliver Wyman Forum (2024)

Na temelju navedenih podataka napravljene su kvote za uzorak primarnog istraživanja, vidljiv u tablici 17.

Tablica 17. Prikaz kvotnog uzorka na 400 ispitanika

	Muškarci	Žene	Ukupno
18 - 24	54	45	99
25 - 34	51	43	94
35 - 44	45	37	82
45 - 54	36	31	67
55 - 64	32	26	58

Ukupno	218	182	400
--------	-----	-----	-----

Izvor: izrada autora

Iz navedenih je podataka vidljivo kako je kod kvotnog uzorka zastupljeniji muški spol te kako su zastupljenije mlađe generacije, a to je u skladu i s drugim istraživanjima na ovo temu, o čemu je ranije u radu bilo riječi (Bick i sur., 2024; Deloitte, 2024; Salesforce, 2024; Lin & Parker, 2025; Aldaroso i sur., 2024a; Aldaroso i sur., 2024b). Prosječna starost ispitanika u ukupnom uzorku je $\approx 37,5$ godina ($\bar{x} = 37,49$, st. dev. = 13,61), pri čemu muškarci čine 54,4 % ispitanika, a žene 45,6 % ispitanika.

7.2.1. Ukupni uzorak

Tablica 18. Demografske karakteristike cjelokupnog uzorka

Kategorije	Broj	Postotak (%)
SPOL		
Muškarci	436	54,4
Žene	366	45,6
DOBNE SKUPINE		
18 – 24	199	24,8
<i>Muškarci</i>	109	13,6
<i>Žene</i>	90	11,2
25 – 34	189	23,6
<i>Muškarci</i>	102	12,7
<i>Žene</i>	87	10,9
35 – 44	164	20,4
<i>Muškarci</i>	90	11,2
<i>Žene</i>	74	9,2
45 – 54	133	16,6
<i>Muškarci</i>	72	9,0
<i>Žene</i>	61	7,6
55 i više	117	14,6
<i>Muškarci</i>	63	7,9
<i>Žene</i>	54	6,7
RADNI STATUS		
Radim	527	65,7

Student sam	141	17,6
Ne radim i nisam student	70	8,7
Student sam, ali i radim	62	7,7
Idem u srednju školu	2	0,3
MJESTO STANOVANJA*		
Ruralno područje	113	14,1
Mali grad	163	20,3
Srednji grad	174	21,7
Veliki grad	352	43,9
RAZINA OBRAZOVANJA		
Osnovna škola	3	0,4
Niža srednja škola (<i>GED / GCSE</i>)	40	5,0
Viša srednja škola (<i>High-school diploma / A-levels</i>)	198	24,7
VSS / Stručni studiji	50	6,2
Prijediplomski studiji	224	27,9
Diplomski studiji	227	28,3
Poslijediplomski studiji	60	7,5

*Kod mjesta stanovanja postoje razlike u definiranju veličine mjesta

Izvor: izrada autora

Ako se zajedno promatraju ispitanici iz Ujedinjenog Kraljevstva i Republike Hrvatske, vidljivo je kako je najkorišteniji alat ChatGPT kojeg koristi čak 95,9 % ispitanika. Slijedi ga Googleov Gemini kojeg koristi 33,9 % ispitanika, a potom slijede Grammarly (29,1 %), Microsoft Copilot (26,2 %), DeepSeek (17, 5%), Perplexity (13,7 %), Dall-E (11,8 %). Quillbot (10,6 %), Midjourney (10,6 %), Grok (8,5 %), Claude (8,4 %), GitHub Copilot (5,4 %), Google NotebookLM (4,6 %), Adobe Firefly (4,4 %), Stable Diffusion (3,5 %), Elicit (3,1 %), SciteAI (2,9 %), Llama (2,5 %), Leonardo (1,7 %), Mistral (1,5 %), JasperAI (1,1 %), Stability (0,9 %), HeyGen (0,7 %), Poe (0,6 %), Runway (0,5 %), AlphaCode (0,2 %) i Ideogram (0,1 %). Neki od dodatno napisanih alata generativne umjetne inteligencije su Harvey, Instatext, ConsensusPro, Blackbox, Pika, Snapchat AI, WordTune, LumaLabs, Suno, Udio, ChatPDF,

Magai, Praktika, Magic School, Meta AI, Talkie AI, Android AI, Canva AI, Chatbox, ElevenLabs, Loveable, Bespoke AI, Athena i Apple Playground AI.

7.3. Rezultati primarnog istraživanja

Za procjenu modela i određivanje ključnih čimbenika koji utječu na prihvatanje i korištenje generativne umjetne inteligencije, prvo je model postavljen u SmartPLS 4 (verzija 4.1.1.2) alatu za statističku obradu podataka (Ringle i sur., 2016). Analize u SmartPLS-u provedene su korištenjem sheme ponderiranja putanja, s postavljenim maksimumom od 3000 iteracija te korištenjem zadane postavke softvera za početne težine. Za procjenu statističke značajnosti rezultata PLS-SEM analize korišten je neparametrijski *bootstrapping* pristup, s ukupno 5000 uzoraka. Ovakav pristup u skladu je s objavljenom literaturom o PLS-SEM analizama (Ringle i sur., 2024; Efron & Tibshirani, 1993; Davison & Hinkley, 1997; Hair i sur., 2020; Hair i sur., 2011).

Znanstvene rasprave dvoje oko uporabe općih pokazatelja prikladnosti modela (engl. *model fit*) u PLS-SEM analizama, jer se često pokazuju neuspješnim i nedovoljnim. Koncept prikladnosti modela definiran u CB-SEM analizama nije primjenjiv za PLS-SEM analize (Hair i sur., 2021). Ipak, PLS modeliranje ima svoje približne kriterije za procjenu prikladnosti modela među kojima je jedini općeprihvaćeni SRMR – standardizirani korijen srednje kvadratne rezidualne pogreške (engl. *Standardized Root Mean Square Residual*) (Hu & Bentler, 1998; Hu & Bentler, 1999; Henseler i sur., 2016). Bryne (2008) navodi kako bi vrijednosti SRMR pokazatelja trebale biti manje od 0,05 kako bi se model smatrao prihvatljivim. Henseler i sur. (2016) pokazali su kako ispravno specificiran model može biti prihvatljiv ako ima vrijednosti od 0,06 ili čak i više. U teoriji se najčešće uzima prag od 0,08 kojeg su predložili Hu i Bentler (1999), a s kojim su se složili Henseler i sur. (2016) te autori SmartPLS softvera Ringle i sur. (2024). Vrijednost SRMR-a koji mjeri prosječnu razliku između empirijskih i modeliranih korelaciju u modelu iznosi 0,050 u cjelovitom modelu (*saturated*), odnosno 0,052 u specificiranom modelu (*estimated*). Obje vrijednosti su ispod 0,08 što ukazuje na dobar *model fit*, drugim riječima model dobro objašnjava podatke (Ringle i sur., 2024).

Kako je vidljivo u tablici 19., vrijednosti u stupcu standardne devijacije u odgovorima po svakoj čestici variraju između 1 i 2 što ukazuje na umjereno do visoko raspršene odgovore, što najčešće znači da tvrdnje dobro diferenciraju odgovore među sudionicima i da postoji dovoljna varijabilnost za statističku analizu. One se mogu tumačiti kao pozitivan pokazatelj, osobito u istraživanju percepcija i stavova. Nadalje, vanjska opterećenja prikazuju apsolutni doprinos

čestice svom pripadajućem konstrukt, a navedene bi vrijednosti trebale prelaziti prag od 0,7, dok se za vrijednosti od 0,4 do 0,7 treba razmotriti eliminacija čestice iz mjernog modela (Hair i sur., 2022). Kako je vidljivo u tablici 19., samo je jedna čestica problematična, a to je FC4. U nastavku ovog potpoglavlja vidljivo je kako je to jedina čestica unutar mjernog modela koja nije zadovoljavajuće razine, ali i kako ona ne narušava značajno diskriminantnu i konvergentnu valjanost niti unutarnju konzistentnost konstrukta. Iz navedenih razloga, autor se odlučio za pristup u kojem čestica FC4 ipak nije eliminirana iz modela radi bolje usporedivosti s rezultatima drugih istraživanja.

U tablici 19. također je vidljivo da je mjerena i VIF vrijednost, odnosno faktor inflacije varijance, kako bi se procijenila kolinearnost među faktorima. Vrijednosti u VIF stupcu pokazuju kako ne postoji značajna kolinearnost među česticama jer niti jedna vrijednost ne prelazi prag od 5,0. Ipak, treba uočiti kako su vrijednosti unutar konstrukta HM relativno visoke, posebice kod čestice HM1 čija je vrijednost 4,895.

Tablica 19. Statistički pokazatelji i validacija mjernih varijabli [Cjeloviti uzorak]

Konstrukt	Prosjek	Standardna devijacija	Vanjsko opterećenje	VIF
Namjera ponašanja				
BI1	5,793	1,204	0,858	2,188
BI2	4,379	1,745	0,866	2,211
BI3	5,192	1,476	0,935	3,304
Očekivani učinak				
PE1	5,521	1,379	0,864	2,276
PE2	5,203	1,433	0,887	2,709
PE3	5,723	1,312	0,879	2,661
PE4	5,292	1,476	0,870	2,405
Očekivani napori				
EE1	5,711	1,228	0,911	3,500
EE2	5,640	1,133	0,862	2,252
EE3	5,779	1,121	0,882	2,734
EE4	5,468	1,223	0,886	2,922
Društveni utjecaj				
SI1	4,118	1,470	0,925	3,218

SI2	4,045	1,493	0,931	3,601
SI3	4,071	1,474	0,938	3,739
Olakšavajući uvjeti				
FC1	5,741	1,187	0,863	2,220
FC2	5,576	1,254	0,862	2,118
FC3	5,590	1,169	0,809	1,649
FC4	5,035	1,456	0,552	1,170
Hedonistička motivacija				
HM1	5,534	1,331	0,952	4,895
HM2	5,539	1,276	0,930	3,352
HM3	5,570	1,318	0,936	4,154
Cijena				
PV1	4,749	1,348	0,938	3,924
PV2	4,708	1,334	0,919	3,073
PV3	4,819	1,354	0,943	4,143
Navika				
HT1	4,221	1,908	0,862	2,236
HT2	2,632	1,710	0,840	2,467
HT3	2,846	1,766	0,803	2,211
HT4	4,265	1,770	0,863	2,175
Osobna inovativnost				
PI1	4,589	1,705	0,928	3,326
PI2	4,133	1,817	0,883	2,258
PI3	5,111	1,577	0,920	3,204
Povjerenje				
TR1	5,887	1,117	0,909	3,679
TR2	5,934	1,087	0,901	3,428
TR3	5,839	1,099	0,883	2,504
TR4	5,789	1,133	0,886	2,747
Korištenje				
USE1	2,520	1,743	0,822	1,768
USE2	4,451	1,779	0,895	2,284
USE3	3,582	1,380	0,903	2,276

Cjeloviti naziv kodiranih čestica pojašnjen u Prilogu 1.

$VIF \leq 5,0$

Vanjsko opterećenje $\geq 0,6$

Tablica 20. prikazuje cjeloviti prikaz svih faktorskih opterećenja u radu, a u kojem je vidljivo kako je čestica FC4 jedina čestica u radu koja ne prelazi prag od 0,70, ali je i vidljivo kako su korelacije između konstrukta EE i konstrukta FC relativno visoke. Ipak, navedena sličnost nije metodološki sporna jer o njima Venkatesh i sur. (2003) pišu još kod predstavljanja izvornog UTAUT modela.

Tablica 20. Unakrsna faktorska opterećenja konstrukta [Cjeloviti uzorak]

	BI	EE	FC	HM	HT	PE	PI	PV	SI	TR	USE
BI1	0,858	0,420	0,413	0,593	0,516	0,624	0,426	0,400	0,358	0,592	0,521
BI2	0,866	0,319	0,285	0,527	0,675	0,549	0,458	0,428	0,459	0,360	0,586
BI3	0,935	0,402	0,384	0,544	0,658	0,604	0,504	0,438	0,402	0,498	0,642
EE1	0,376	0,911	0,658	0,432	0,340	0,457	0,371	0,272	0,250	0,331	0,347
EE2	0,403	0,862	0,589	0,511	0,363	0,481	0,353	0,331	0,301	0,330	0,397
EE3	0,356	0,882	0,607	0,443	0,301	0,407	0,264	0,275	0,257	0,321	0,311
EE4	0,379	0,886	0,637	0,435	0,344	0,452	0,352	0,267	0,294	0,318	0,359
FC1	0,323	0,589	0,863	0,326	0,264	0,331	0,310	0,235	0,183	0,349	0,323
FC2	0,337	0,702	0,862	0,344	0,308	0,345	0,340	0,259	0,160	0,286	0,360
FC3	0,367	0,521	0,809	0,325	0,315	0,369	0,354	0,296	0,226	0,401	0,325
FC4	0,224	0,336	0,552	0,305	0,242	0,225	0,143	0,226	0,267	0,166	0,196
HM1	0,584	0,493	0,389	0,952	0,451	0,577	0,392	0,410	0,317	0,436	0,423
HM2	0,615	0,521	0,427	0,930	0,455	0,591	0,416	0,422	0,340	0,462	0,465
HM3	0,557	0,433	0,339	0,936	0,429	0,552	0,391	0,416	0,338	0,427	0,404
HT1	0,659	0,367	0,321	0,494	0,862	0,601	0,392	0,461	0,433	0,309	0,637
HT2	0,490	0,221	0,224	0,299	0,840	0,388	0,340	0,329	0,339	0,184	0,512
HT3	0,495	0,220	0,216	0,280	0,803	0,355	0,336	0,356	0,387	0,201	0,456
HT4	0,663	0,431	0,417	0,475	0,863	0,567	0,510	0,384	0,383	0,374	0,654
PE1	0,609	0,440	0,333	0,570	0,528	0,864	0,343	0,411	0,407	0,440	0,523
PE2	0,583	0,438	0,347	0,533	0,542	0,887	0,339	0,430	0,406	0,394	0,527
PE3	0,542	0,473	0,417	0,525	0,497	0,879	0,348	0,377	0,349	0,436	0,505
PE4	0,594	0,432	0,351	0,510	0,473	0,870	0,376	0,396	0,410	0,430	0,493
PI1	0,490	0,331	0,337	0,383	0,442	0,362	0,928	0,251	0,228	0,330	0,468
PI2	0,465	0,341	0,331	0,351	0,451	0,377	0,883	0,233	0,283	0,293	0,497
PI3	0,472	0,367	0,372	0,429	0,409	0,359	0,920	0,237	0,211	0,378	0,469
PV1	0,444	0,291	0,290	0,416	0,422	0,424	0,235	0,938	0,276	0,323	0,332
PV2	0,440	0,294	0,292	0,388	0,420	0,417	0,270	0,919	0,266	0,302	0,357
PV3	0,448	0,324	0,325	0,436	0,444	0,453	0,234	0,943	0,300	0,358	0,356

SI1	0,429	0,320	0,270	0,336	0,423	0,428	0,267	0,284	0,925	0,261	0,344
SI2	0,411	0,247	0,224	0,306	0,418	0,389	0,226	0,258	0,931	0,237	0,332
SI3	0,439	0,304	0,219	0,342	0,443	0,440	0,243	0,298	0,938	0,237	0,365
TR1	0,459	0,315	0,365	0,391	0,272	0,413	0,316	0,291	0,216	0,909	0,311
TR2	0,462	0,315	0,349	0,425	0,272	0,442	0,328	0,293	0,227	0,901	0,313
TR3	0,536	0,350	0,359	0,434	0,337	0,456	0,362	0,361	0,277	0,883	0,371
TR4	0,479	0,331	0,344	0,432	0,288	0,425	0,302	0,303	0,215	0,886	0,319
USE1	0,467	0,229	0,272	0,313	0,510	0,412	0,383	0,266	0,255	0,243	0,822
USE2	0,625	0,355	0,337	0,432	0,590	0,539	0,465	0,319	0,365	0,387	0,895
USE3	0,619	0,442	0,408	0,446	0,674	0,567	0,514	0,381	0,346	0,326	0,903

Pokazatelji unutarnje konzistentnosti te konvergentne valjanosti na temelju konstrukta prikazane su u tablici 21. Vrijednosti Cronbachove alfe i kompozitne pouzdanosti trebali bi prelaziti 0,7 kako bi se smatrale valjanima, a kako je vidljivo u tablici 21. niti jedna vrijednost nije ispod 0,7 što znači da valjanost niti jednog konstrukta nije narušena iako je konstrukt olakšavajućih uvjeta svugdje najslabiji ($\alpha = 0,779$; CR (rho_c) = 0,860; CR (rho_a) = 0,815). Slična je stvar i s pokazateljem AVE koji ne bi trebao biti ispod 0,5. Kako je vidljivo u tablici 21., konvergentna valjanost niti jednog konstrukta nije narušena, no ponovno je vrijednost konstrukta FC (AVE = 0,612) značajno niža nego kod drugih konstrukta.

Tablica 21. Pokazatelji unutarnje konzistentnosti i konvergentne valjanosti [Cjeloviti uzorak]

Konstrukt	Cronbach alfa (α)	CR (rho_c)	CR (rho_a)	AVE
Namjera ponašanja (BI)	0,864	0,917	0,869	0,787
Očekivani učinak (PE)	0,898	0,929	0,899	0,766
Očekivani naponi (EE)	0,908	0,936	0,909	0,784
Društveni utjecaj (SI)	0,924	0,952	0,925	0,868
Olakšavajući uvjeti (FC)	0,779	0,860	0,815	0,612
Hedonistička motivacija (HM)	0,933	0,957	0,935	0,882
Cijena (PV)	0,926	0,953	0,926	0,871
Navika (HT)	0,865	0,907	0,880	0,709
Osobna inovativnost (PI)	0,897	0,936	0,898	0,829
Povjerenje (TR)	0,917	0,942	0,920	0,801
Korištenje (USE)	0,846	0,906	0,862	0,764

AVE $\geq 0,5$

CR (rho_a; rho_c), Cronbach's $\geq 0,7$

U svrhu procjene diskriminantne valjanosti koriste se Fornell Larckerov kriterij i HTMT odnos. U tablici 22. vidljivo je kako se u Fornell-Larckerovom kriteriju korelacije među latentnim varijablama uspoređuju sa korijenom pokazatelja AVE svake od njih. S obzirom na to da kriterij nalaže kako bi korijen pokazatelja AVE trebao biti veći od bilo koje korelacije tog konstrukta s nekim drugim konstruktom, vidljivo je kako diskriminantna valjanost konstrukta nije narušena. Ipak, treba naglasiti kako je vidljiva visoka korelacija između konstrukta EE i konstrukta FC koja iznosi čak 0,703. Iako navedena korelacija ne prelazi granicu, odnosno nije veća od korijena pokazatelja AVE, vidljivo je da između navedenih faktora postoji jako velika povezanost što potvrđuje njihovu, ranije spomenutu, povezanost (Shaw & Sergueeva, 2019;

Venkatesh i sur., 2003). Također, vidljiva je i visoka korelacija između konstrukta HM i konstrukta PE, što dokazuje snažnu povezanost između korištenja radi funkcionalnosti alata generativne umjetne inteligencije i užitka koje pruža to korištenje. Ipak, navedena korelacija ne narušava diskriminantnu valjanost mjernog modela.

Tablica 22. Diskriminantna valjanost prema Fornell-Larckerovom kriteriju
[Cjeloviti uzorak]

	BI	EE	FC	HM	HT	PE	PI	PV	SI	TR	USE
BI	0,887										
EE	0,428	0,886									
FC	0,406	0,703	0,782								
HM	0,624	0,516	0,412	0,939							
HT	0,697	0,382	0,362	0,474	0,842						
PE	0,667	0,509	0,412	0,611	0,583	0,875					
PI	0,523	0,380	0,381	0,426	0,477	0,402	0,911				
PV	0,476	0,324	0,324	0,443	0,459	0,462	0,264	0,933			
SI	0,458	0,312	0,255	0,353	0,460	0,450	0,264	0,301	0,931		
TR	0,543	0,368	0,396	0,471	0,329	0,486	0,367	0,351	0,263	0,895	
USE	0,659	0,401	0,393	0,460	0,683	0,586	0,525	0,373	0,373	0,369	0,874

Heterotrait-Monotrait ratio ili jednostavnije HTMT kriterij, suvremeni je pristup mjerenja diskriminantne valjanosti koji mjeri prosječnu sličnost između čestica koje bi trebale biti neki drugi konstrukt s česticama koje bi trebale pripadati istom konstrukt. Ukoliko bi vrijednost HTMT omjera prelazila prag od 0,90 to bi označavalo narušenost mjernog modela i preveliku sličnost među česticama konstrukta, no vidljivo je kako jedna vrijednost ne prelazi taj prag premda je vrijednost od 0,823 između konstrukta EE i FC vrlo visoka, kao i kod Fornell-Larckerovog kriterija. Podatci mjernog modela ovime su se pokazali pouzdanima, što pruža čvrstu osnovu za daljnju validaciju i interpretaciju mjernog modela.

Tablica 23. Procjena diskriminantne valjanosti korištenjem HTMT kriterija
[Cjeloviti uzorak]

	BI	EE	FC	HM	HT	PE	PI	PV	SI	TR
EE	0,483									
FC	0,492	0,823								
HM	0,696	0,556	0,489							
HT	0,789	0,413	0,429	0,510						
PE	0,757	0,563	0,491	0,666	0,642					
PI	0,593	0,420	0,442	0,465	0,531	0,448				
PV	0,532	0,352	0,385	0,476	0,507	0,506	0,289			
SI	0,512	0,339	0,317	0,380	0,511	0,492	0,290	0,325		
TR	0,609	0,401	0,458	0,507	0,353	0,534	0,402	0,378	0,284	
USE	0,761	0,444	0,471	0,510	0,776	0,664	0,597	0,416	0,417	0,412

HTMT $\leq$ 0,90

7.3.1. Testiranje hipoteza

Za testiranje hipoteza unutar istraživanja korišten je *bootstrapping* kako bi se utvrdila značajnost procijenjenih koeficijenata putanje unutar PLS-SEM modeliranja. Kako je i prethodno navedeno, korištena je metoda od preporučenih 5000 *bootstrap* uzoraka. Iz tablice 24. vidljivo je kako se moderatori nisu pokazali osobito značajnima. Unutar konceptualnog modela samo je jedan moderator utjecao na odnos između zavisne i nezavisne varijable, a to je utjecaj spola na odnos između hedonističke motivacije i namjere ponašanja.

Kako je i u metodološkom dijelu rečeno, u odluci o prihvatanju ili odbacivanju hipoteze primarno se promatra odnos između nezavisne i zavisne varijable, a potom sekundarno utjecaj moderatora na taj odnos. U nastavku su pojašnjene odluke o prihvatanju i odbacivanju hipoteza temelje na rezultatima unutar tablice 24. koja prikazuje rezultate koeficijenata putanje, t-vrijednosti i p-vrijednosti za svaki odnos unutar hipoteza.

- **H1:** *Dob i spol moderiraju utjecaj očekivanog učinka na namjeru ponašanja, tako da je utjecaj jači kod mlađih muškaraca.*

Očekivani učinak kod korištenja alata generativne umjetne inteligencije statistički značajno utječe na namjeru ponašanja, odnosno na namjeru korisnika da u budućnosti koristi navedene alate ($\beta = 0,113$; $t = 2,532$; $p < 0,05$). S druge strane, vidljivo je i kako dob ($\beta = 0,015$; $t = 0,462$; $p > 0,05$) i spol ($\beta = 0,135$; $t = 1,774$; $p > 0,05$) nisu statistički značajni moderatori u tom odnosu. Premda su t-vrijednost i p-vrijednost kod spola blizu praga značajnosti, oni nisu statistički značajni. S obzirom na to da se izravni utjecaj očekivanog učinka (PE) na namjeru ponašanja (BI) pokazao statistički značajnim, odluka je kako se hipoteza 1 djelomično prihvata.

- **H2:** *Dob, spol i iskustvo moderiraju utjecaj očekivanih napora na namjeru ponašanja, tako da je utjecaj jači kod mlađih žena u ranoj fazi iskustva s tehnologijom.*

Očekivani napori kod korištenja alata generativne umjetne inteligencije ne utječu statistički značajno na namjeru ponašanja, odnosno na namjeru korisnika o budućem korištenju takve tehnologije ($\beta = -0,021$; $t = 0,457$; $p > 0,05$). Jednako tako, ni dob ($\beta = -0,025$; $t = 0,690$; $p > 0,05$), ni spol ($\beta = -0,015$; $t = 0,207$; $p > 0,05$) niti iskustvo ($\beta = 0,011$; $t = 0,278$; $p > 0,05$) nisu se istaknuli kao statistički značajni moderatori te veze. Budući da ni izravni utjecaj očekivanih napora (EE) na namjeru ponašanja (BI) nije statistički značajan, niti je potvrđena značajnost moderatora navedenog odnosa, odluka je kako se hipoteza 2 odbacuje.

- **H3:** *Dob, spol i iskustvo moderiraju utjecaj društvenog utjecaja na namjeru ponašanja, tako da je utjecaj jači kod starijih žena u ranoj fazi iskustva s tehnologijom.*

Konstrukt društveni utjecaja pokazao se kao statistički značajan u utjecaju na namjeru ponašanja ($\beta = 0,075$; $t = 2,074$; $p < 0,05$). Nasuprot izravnom utjecaju nezavisne varijable, moderatori dob ($\beta = -0,004$; $t = 0,130$; $p > 0,05$), spol ($\beta = -0,028$; $t = 0,535$; $p > 0,05$) i iskustvo ($\beta = 0,020$; $t = 0,663$; $p > 0,05$) nisu se potvrdili kao statistički značajni. Unatoč nedostatku potvrđenih moderacijskih efekata, budući da se izravni utjecaj SI na BI pokazao statistički značajnim, odluka je kako se hipoteza 3 djelomično prihvaća.

- **H4a:** *Dob, spol i iskustvo moderiraju utjecaj olakšavajućih uvjeta na namjeru ponašanja, tako da je utjecaj jači kod starijih žena u ranoj fazi iskustva s tehnologijom.*

S obzirom na to da se u konceptualnom modelu ispituje utjecaj olakšavajućih uvjeta i na zavisnu varijablu namjere ponašanja, ali i na zavisnu varijablu korištenja, navedena je hipoteza podijeljena na dva dijela.

U prvom dijelu hipoteze koji ispituje utjecaj olakšavajućih uvjeta na namjeru ponašanja, potvrđeno je kako on nije statistički značajan ($\beta = 0,046$; $t = 1,070$; $p > 0,05$). Nadalje, pokazalo se kako dob ($\beta = 0,047$; $t = 1,330$; $p > 0,05$), spol ($\beta = -0,053$; $t = 0,829$; $p > 0,05$) i iskustvo ($\beta = 0,034$; $t = 0,863$; $p > 0,05$) kao moderatori navedenog odnosa također nisu statistički značajni. S obzirom na to da niti utjecaj FC na BI, a niti moderatori nisu potvrđeni kao statistički značajni, hipoteza 4a se odbacuje.

- **H4b:** *Dob i iskustvo moderiraju utjecaj olakšavajućih uvjeta na korištenje, tako da je utjecaj jači kod starijih, posebno onih s većom razinom iskustva s tehnologijom.*

U drugom dijelu hipoteze koji ispituje utjecaj olakšavajućih uvjeta na korištenje alata generativne umjetne inteligencije, navedeni je utjecaj potvrđen kao statistički značajan ($\beta = 0,086$; $t = 3,105$; $p < 0,05$). Ipak, moderatori dob ($\beta = -0,001$; $t = 0,041$; $p > 0,05$) i iskustvo ($\beta = 0,028$; $t = 1,160$; $p > 0,05$) nisu potvrđeni kao statistički značajni. Sukladno ovim nalazima koji potvrđuju glavni efekt, ali ne i moderacijske učinke, hipoteza H4b djelomično se prihvaća.

- **H5:** *Dob, spol i iskustvo moderiraju utjecaj hedonističke motivacije na namjeru ponašanja, tako da je utjecaj jači među mlađim muškarcima u ranoj fazi iskustva s tehnologijom.*

Utjecaj hedonističke motivacije pokazao se statistički značajnim i snažnim prediktorom namjere ponašanja ($\beta = 0,239$; $t = 5,452$; $p < 0,05$). Također, spol se pokazao kao statistički

značajan moderator navedenog odnosa ($\beta = -0,148$; $t = 2,053$; $p < 0,05$), što ukazuje na to da je utjecaj snažniji kod muškaraca, s druge strane dob ($\beta = -0,034$; $t = 1,103$; $p > 0,05$) i iskustvo ($\beta = -0,037$; $t = 1,081$; $p > 0,05$) se nisu pokazali statistički značajnim moderatorima. Shodno navedenim nalazima koji potvrđuju glavni utjecaj i jedan značajan moderacijski učinak, odluka je kako se hipoteza 5 djelomično prihvaća.

- **H6:** *Dob i spol moderiraju utjecaj vrijednosti cijene na namjeru ponašanja, tako da je utjecaj jači među starijim ženama.*

Cijena, odnosno konstrukt vrijednosti cijene nije statistički značajan prediktor namjere ponašanja ($\beta = 0,023$; $t = 0,700$; $p > 0,05$). Također, moderatori dob ($\beta = -0,028$; $t = 0,965$; $p > 0,05$) i spol ($\beta = 0,054$; $t = 0,912$; $p > 0,05$) ne moderiraju navedeni odnos u značajnoj mjeri. Slijedom ovih rezultata, hipoteza 6 se odbacuje.

- **H7a:** *Dob, spol i iskustvo moderiraju utjecaj navike na namjeru ponašanja, tako da je utjecaj jači kod starijih muškaraca s većom razinom iskustva s tehnologijom.*

Navika, baš kao i olakšavajući uvjeti, nezavisna je varijabla čiji se utjecaj unutar ovog konceptualnog modela ispituje u dva smjera, na zavisnu varijablu namjere ponašanja, ali i na zavisnu varijablu korištenja te je zbog toga ova hipoteza podijeljena u dva dijela.

Prvi dio hipoteze koji ispituje utjecaj navike na namjeru ponašanja potvrđuje kako je navika naj snažniji prediktor namjere ponašanja u ovom istraživanju te kako ima statistički značajan utjecaj na namjeru ponašanja ($\beta = 0,336$; $t = 8,972$; $p < 0,05$). Ipak, rezultati pokazuju kako dob ($\beta = -0,009$; $t = 0,299$; $p > 0,05$), spol ($\beta = 0,052$; $t = 0,939$; $p > 0,05$) i iskustvo ($\beta = -0,022$; $t = 0,799$; $p > 0,05$) nisu značajni moderatori navedenog utjecaja. U skladu s navedenim rezultatima, odluka je kako se hipoteza 7a djelomično prihvaća.

- **H7b:** *Dob, spol i iskustvo moderiraju utjecaj navike na korištenje tehnologije, tako da je utjecaj jači kod starijih muškaraca s većom razinom iskustva s tehnologijom.*

Drugi dio hipoteze koji ispituje utjecaj navike na korištenje alata generativne umjetne inteligencije također se pokazao statistički značajnim, ali i također naj snažnijim prediktorom korištenja navedene tehnologije ($\beta = 0,409$; $t = 10,030$; $p < 0,05$). Ipak, rezultati pokazuju kako dob ($\beta = 0,034$; $t = 1,276$; $p > 0,05$), spol ($\beta = -0,042$; $t = 0,920$; $p > 0,05$) i iskustvo ($\beta = -0,001$; $t = 0,025$; $p > 0,05$) nisu statistički značajni moderatori navedenog odnosa. U skladu s navedenim rezultatima, zaključak je kako se hipoteza 7b djelomično prihvaća.

- **H8:** *Dob i spol moderiraju utjecaj osobne inovativnosti na namjeru ponašanja, tako da je utjecaj jači kod mlađih muškaraca.*

Konstrukat osobne inovativnosti, kojim je izvorni UTAUT2 proširen u ovom istraživanju, statistički je značajan prediktor namjere ponašanja ($\beta = 0,070$; $t = 1,963$; $p \leq 0,05$). Iako je riječ o vrlo graničnoj vrijednosti jer vidljivo je kako je t-vrijednost tek nešto iznad 1,96 te kako je p-vrijednost točno 0,050 što po teoriji ipak ulazi u domenu rubno prihvatljivog, ali i dalje statistički značajnog rezultata (Rumsey, 2016; Madjarova i sur., 2022; Greenland i sur., 2016). Ipak, moderatori dob ($\beta = -0,012$; $t = 0,466$; $p > 0,05$) i spol ($\beta = 0,065$; $t = 1,285$; $p > 0,05$) nisu statistički značajni. Temeljem ovih nalaza, odluka je kako se hipoteza 8 djelomično prihvaća.

- **H9:** *Dob, spol i iskustvo moderiraju utjecaj povjerenja na namjeru ponašanja, tako da je utjecaj jači među mlađim muškarcima s većom razinom iskustva s tehnologijom.*

Konstrukat povjerenja još je jedan od konstrukta kojim je izvorni UTAUT2 proširen u sklopu ovog istraživanja, a tablica 24 potvrđuje kako povjerenje statistički značajno utječe na namjeru ponašanja ($\beta = 0,187$; $t = 5,424$; $p < 0,05$). Ipak, moderatori dob ($\beta = 0,018$; $t = 0,657$; $p > 0,05$), spol ($\beta = 0,005$; $t = 0,095$; $p > 0,05$) i iskustvo ($\beta = -0,037$; $t = 1,285$; $p > 0,05$) nisu se pokazali statistički značajnima moderatorima navedenog odnosa. Sukladno potvrđenom izravnom utjecaju povjerenja na namjeru ponašanja, odluka je kako se hipoteza 9 djelomično prihvaća.

- **H10:** *Iskustvo moderira utjecaj namjere ponašanja na korištenje tehnologije, tako da je utjecaj jači kod korisnika s manje iskustva.*

U konačnici, u sklopu ovog konceptualnog modela ispituje se i izravni utjecaj namjere ponašanja na stvarno korištenje te se isto pokazalo statistički značajno ($\beta = 0,310$; $t = 9,123$; $p < 0,05$). Ipak, iskustvo se nije pokazalo kao statistički značajan moderator ovog odnosa ($\beta = 0,039$; $t = 1,339$; $p > 0,05$). Shodno ovim rezultatima, odluka je kako se hipoteza 10 djelomično prihvaća.

Tablica 24. Rezultati testiranja hipoteza istraživanja [Cjeloviti uzorak]

Hipoteza	Odnos	β	t-vrijednost	p	Odluka
H1	PE > BI	0,113	2,532	0,011*	Hipoteza se djelomično prihvaća.
	Dob x PE > BI	0,015	0,462	0,644	
	Spol x PE > BI	0,135	1,774	0,076	
H2	EE > BI	-0,021	0,457	0,648	Hipoteza se odbacuje.
	Dob x EE > BI	-0,025	0,690	0,490	
	Spol x EE > BI	-0,015	0,207	0,836	

	Iskustvo x EE > BI	0,011	0,278	0,781	
H3	SI > BI	0,075	2,074	0,038*	Hipoteza se djelomično prihvaća.
	Dob x SI > BI	-0,004	0,130	0,897	
	Spol x SI > BI	-0,028	0,535	0,592	
	Iskustvo x SI > BI	0,020	0,663	0,507	
H4a	FC > BI	0,046	1,070	0,285	Hipoteza se odbacuje.
	Dob x FC > BI	0,047	1,330	0,184	
	Spol x FC > BI	-0,053	0,829	0,407	
	Iskustvo x FC > BI	0,034	0,863	0,388	
H4b	FC > USE	0,086	3,105	0,002**	Hipoteza se djelomično prihvaća.
	Dob x FC > USE	-0,001	0,041	0,967	
	Iskustvo x FC > USE	0,028	1,160	0,246	
H5	HM > BI	0,239	5,452	0,000***	Hipoteza se djelomično prihvaća.
	Dob x HM > BI	-0,034	1,103	0,270	
	Spol x HM > BI	-0,148	2,053	0,040*	
	Iskustvo x HM > BI	-0,037	1,081	0,280	
H6	PV > BI	0,023	0,700	0,484	Hipoteza se odbacuje.
	Dob x PV > USE	-0,028	0,965	0,334	
	Spol x PV > USE	0,054	0,912	0,362	
H7a	HT > BI	0,336	8,972	0,000***	Hipoteza se djelomično prihvaća.
	Dob x HT > BI	-0,009	0,299	0,765	
	Spol x HT > BI	0,052	0,939	0,348	
	Iskustvo x HT > BI	-0,022	0,799	0,424	
H7b	HT > USE	0,409	10,030	0,000***	Hipoteza se djelomično prihvaća.
	Dob x HT > USE	0,034	1,276	0,202	
	Spol x HT > USE	-0,042	0,920	0,358	
	Iskustvo x HT > USE	-0,001	0,025	0,980	
H8	PI > BI	0,070	1,963	0,050*	Hipoteza se djelomično prihvaća.
	Dob x PI > BI	-0,012	0,466	0,641	
	Spol x PI > BI	0,065	1,226	0,220	
H9	TR > BI	0,187	5,424	0,000***	Hipoteza se djelomično prihvaća.
	Dob x TR > BI	0,018	0,657	0,511	
	Spol x TR > BI	0,005	0,095	0,925	
	Iskustvo x TR > BI	-0,037	1,285	0,199	
H10	BI > USE	0,310	9,123	0,000***	Hipoteza se djelomično prihvaća.
	Iskustvo x BI > USE	0,039	1,339	0,181	

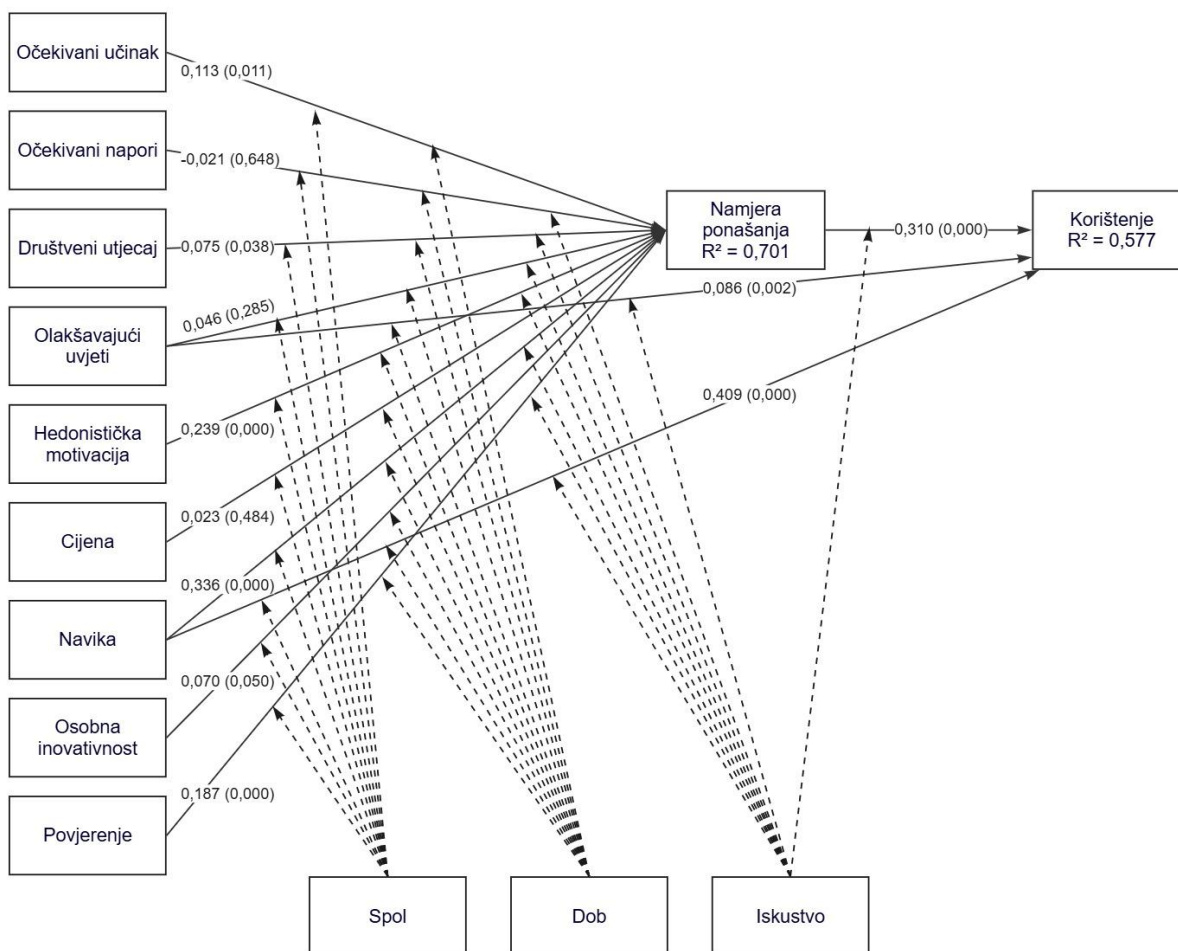
β (Beta) = *Path coefficients* (standardizirani koeficijent putova); p = p vrijednost

* $p \leq 0,05$; ** $p < 0,01$; *** $p < 0,001$

Osim spomenutih hipoteza, u konceptualnom dizajnu istraživanja postavljeno je i istraživačko pitanje koje objedinjuje sve najvažnije odnose u konceptualnom modelu, a ono glasi: *Je li i u kojoj mjeri moguće predvidjeti namjeru ponašanja i stvarno ponašanje korisnika generativne umjetne inteligencije na temelju predloženog konceptualnog modela koji je utemeljen na UTAUT2 modelu (očekivani učinak, očekivani naponi, društveni utjecaj, olakšavajući uvjeti,*

hedonistička motivacija, cijena i navika moderirani dobi, spolom i iskustvom) te proširen konstruktima osobne inovativnosti i povjerenjem?

Sukladno vrijednostima navedenim u tablici 24. i prethodno navedenoj analizi hipoteza koja navodi kako se 9 od ukupno 12 pretpostavki djelomično prihvaća te uz nadopunu kako je u sklopu ovog istraživanja analizirano i objašnjenje varijance modela koja kaže kako je R^2 za namjeru ponašanja 0,701 (odnosno 70,1 % objašnjenja varijance), a R^2 za stvarno korištenje 0,577 (odnosno 57,7 % objašnjenja varijance). Zaključno, odgovor na istraživačko pitanje glasi – proširenim UTAUT2 modelom moguće je predvidjeti namjeru ponašanja i stvarno korištenje generativne umjerene inteligencije, uz objašnjenje modela od 70,1 % kod namjere ponašanja i 57,7 % kod stvarnog korištenja. Pojednostavljeni prikaz vidljiv je na slici 18. Radi urednijeg prikaza modela, na putanje moderatora nisu upisane vrijednosti, one su prikazane u tablici 24. S druge strane, osim R^2 vrijednosti kao koeficijenta determinacije, prijavljuje se i Q^2 ili Stone-Geisserov koeficijent predikcije (Stone, 1974; Geisser, 1974; Hair i sur., 2022). Q^2 mjeri koliko dobro model može predvidjeti izostavljene podatke za zavisne varijable, drugim riječima, on mjeri prediktivnu relevantnost modela. Q^2 za namjeru ponašanja (BI) iznosi 0,659 dok je Q^2 za zavisnu varijablu korištenja 0,546. Oba rezultata imaju vrijednost veću od 0,50 što upućuje na izrazito snažnu prediktivnu relevantnost modela. S obzirom na visoku sličnosti između vrijednosti R^2 i Q^2 za zavisnu varijablu BI i zavisnu varijablu USE, može se zaključiti kako model pokazuje visoku stabilnost predikcije i izbjegava problem pretjeranog prilagođavanja podacima.



* koeficijent putanje - β (vrijednost ispred zagrade); p-vrijednost (vrijednost u zagradi)

Slika 18. Rezultati empirijskog istraživanja na konceptualnom modelu [Cjeloviti uzorak]

Izvor: izrada autora

Ukoliko bi se promatrao model bez moderatora, SRMR vrijednost za cjeloviti model iznosi 0,053, dok za specificirani iznosi 0,054. Obje vrijednosti zadovoljavajuće su razine ($< 0,080$). Objašnjenje varijance kod namjere ponašanja je 67,9 % ($R^2 = 0,679$) dok je kod korištenja 54,2 % ($R^2 = 0,542$). Stone-Geisserov koeficijent predikcije za namjeru ponašanja iznosi 0,668, dok za korištenje iznosi 0,525. Ponovno je vidljivo kako su R^2 vrijednosti i Q^2 vrijednosti približno slične što upućuje na visoku stabilnost predikcije modela.

Dakle, u modelu bez moderatora objašnjenje varijance niže je za 2,2 postotna poena kod BI i 3,5 postotnih poena kod USE. Kako je vidljivo u tablici 25., ispod teksta, svi prediktori koji su potvrđeni kao statistički značajni u analizi na modelu s moderatorima, potvrđeni su i ovdje. Ipak, treba napomenuti kako je vidljiv porast utjecaja $PE > BI$ i $PI > BI$ kada na njih nisu stavljeni moderirajući efekti. S druge strane, $EE > BI$, $FC > BI$ i $PV > BI$ koji nisu potvrđeni u analizi konceptualnog modela, nisu potvrđeni niti ovdje.

Tablica 25. Testiranje odnosa u modelu bez moderatora [Cjeloviti uzorak]

Odnos	β	t-vrijednost	p
PE > BI	0,178	5,199	0,000***
EE > BI	-0,052	1,402	0,161
SI > BI	0,070	2,575	0,010**
FC > BI	0,013	0,377	0,706
FC > USE	0,111	4,168	0,000***
HM > BI	0,187	5,281	0,000***
PV > BI	0,045	1,581	0,114
HT > BI	0,347	12,791	0,000***
HT > USE	0,417	12,251	0,000***
PI > BI	0,121	4,523	0,000***
TR > BI	0,190	6,933	0,000***
BI > USE	0,324	9,599	0,000***

β (Beta) = *Path coefficients* (standardizirani koeficijent putova); p = p vrijednost

* $p \leq 0,05$; ** $p < 0,01$; *** $p < 0,001$

7.4. Analiza rezultata s poduzorka iz Ujedinjenog Kraljevstva

7.4.1. Prikupljanje podataka i uzorak na tržištu Ujedinjenog Kraljevstva

U procesu regrutiranja ispitanika na tržištu Ujedinjenog Kraljevstva korištena je platforma Prolific (Prolific, 2025). U procesu filtriranja željenih ispitanika odabrani su ispitanici koji žive na području Ujedinjenog Kraljevstva, znaju engleski jezik te koji su prethodno imali iskustva s korištenjem nekog od alata generativne umjetne inteligencije. Ostale kategorije napravljene su shodno ranije definiranim kvotama. Početak provedbe prikupljanja podataka krenuo je 4. ožujka 2025. godine i trajao je do 13. ožujka 2025. godine. Prikupljeno je ukupno 420 ispitanika, od čega je pet ispitanika isključeno zbog loše kvalitete podataka, po prijedlogu Alchemerovog algoritma koji isključuje odgovore ispitanika ukoliko su u 1 % najbrže ispunjenih, imaju previše sličan obrazac ispunjavanja ili su nedosljedni u ispunjavanju. Nadalje, 11 ispitanika nije dovršilo svoje ispunjavanje, a sedam ispitanika diskvalificirano je na prvom pitanju koje ispituje jesu li ispitanici ikada koristili neki od alata generativne umjetne inteligencije. U konačnici, ukupno 402 ispitanika dovršili su proces ispunjavanja istraživačkog upitnika sa zadovoljavajućom razinom kvalitete podataka.

Tablica 26. Demografske karakteristike uzorka na tržištu Ujedinjenog Kraljevstva

Kategorije	Broj	Postotak (%)
SPOL		
Muškarci	218	54,2
Žene	184	45,8
DOBNE SKUPINE		
18 – 24	100	24,9
<i>Muškarci</i>	55	13,7
<i>Žene</i>	45	11,2
25 – 34	95	23,6
<i>Muškarci</i>	51	12,7
<i>Žene</i>	44	10,9
35 – 44	82	20,4
<i>Muškarci</i>	45	11,2
<i>Žene</i>	37	9,2
45 – 54	66	16,4
<i>Muškarci</i>	36	9,0

<i>Žene</i>	30	7,5
55 – 65	59	14,7
<i>Muškarci</i>	31	7,7
<i>Žene</i>	28	7,0
RADNI STATUS		
Radim	268	66,7
Student sam	58	14,4
Ne radim i nisam student	58	14,4
Student sam, ali i radim	16	4,0
Idem u srednju školu	2	0,5
MJESTO STANOVANJA		
Ruralno područje (<i>manje od 10,000 stanovnika</i>)	53	13,2
Mali grad (<i>10,000 – 75,000 stanovnika</i>)	119	29,6
Srednji grad (<i>75,000 – 250,000 stanovnika</i>)	106	26,4
Veliki grad (<i>više od 250,000 stanovnika</i>)	124	30,8
RAZINA OBRAZOVANJA		
Osnovna škola	3	0,7
Niža srednja škola (<i>GED / GCSE</i>)	40	10,0
Viša srednja škola (<i>High-school diploma / A-levels</i>)	78	19,4
VSS / Stručni studiji	38	9,5
Prijediplomski studiji	140	34,8
Diplomski studiji	89	22,1
Poslijediplomski studiji	14	3,5

Izvor: izrada autora prema rezultatima primarnog istraživanja

Najčešće korišteni alat među ispitanicima s tržišta Ujedinjenog Kraljevstva je ChatGPT, a kojeg koristi 94,8 % ispitanika, slijede ga Gemini (41,3 %), Microsoft Copilot (33,8 %), Grammarly (28,6 %), DeepSeek (16,4 %), Dall-E (14,9 %), Grok (12,7 %), Quillbot (11,7 %), Midjourney (10,7 %), Claude (10,2 %), GitHub Copilot (6,5 %), Google Notebook (5,5 %), Perplexity (5,5 %), Adobe Firefly (5 %), Stable Diffusion (4 %), Meta Llama (3 %), Mistral (1,7 %), Leonardo (1,5 %), Stability AI (1 %), Poe (1 %), Runway (0,5 %), Elicit (0,5 %), HeyGen (0,5 %), Scite AI (0,5 %), Jasper AI (0,5 %), Ideogram (0,3 %), a od nadopisanih odgovora još se spominju i

Meta AI, Snapchat AI, Character.ai, Athena, Canva AI, Android AI, Loveable, Bespoke AI, Talkie AI, Magicschool, Chatbox, Remini, Qwen, Apple Playground AI i ElevenLabs.

7.4.2. Rezultati primarnog istraživanja na tržištu Ujedinjenog kraljevstva

Kako je prethodno spomenuto, kod PLS-SEM analize za procjenu prikladnosti modela koristi se SRMR čija bi vrijednost morala biti ispod 0,08 kako bi se model smatrao prihvatljivim. U slučaju analize ispitanika s područja Ujedinjenog Kraljevstva, SRMR vrijednost kod cjelovitog modela (*saturated*) iznosi 0,054, a kod specificiranog modela (*estimated*) navedena vrijednost iznosi 0,056. Obje vrijednosti su ispod praga od 0,08 što ukazuje na dobar *model fit*, odnosno model dobro objašnjava podatke.

Tablica 27. pokazuje prosječne vrijednosti i standardnu devijaciju za svaku česticu unutar mjernog modela, ali i vanjska opterećenja i faktor inflacije varijance. Zanimljivo je za promotriti kako su vrijednosti čestica unutar konstrukta osobne inovativnosti vrlo niske u odnosu na čestice iz drugih konstrukta, ali i u usporedbi s uzorkom s tržišta Republike Hrvatske. Gotovo sve vrijednosti standardne devijacije su između 1 i 2 što ukazuje na umjereno do visoko raspršene odgovore. Najmanje standardne devijacije vidljive su kod čestica unutar konstrukta društvenog utjecaja. Kod stupca vanjskih opterećenja vidljivo je kako je jedna vrijednost ispod praga od 0,60. Riječ je o čestici FC4 koja je prethodno protumačena u analizi podataka na bazi cjelokupnog uzorka. Ipak, ono što nije bio slučaj na cjelovitom uzorku jesu povišene VIF vrijednosti kod čestica HM1 i HM2 koje prelaze predloženi prag od 5,0. Kako je vidljivo u Prilogu 1, riječ je o česticama:

HM1 – Korištenje generativne umjetne inteligencije je zabavno (*engl. Using Gen AI is fun*).

HM2 – Korištenje generativne umjetne inteligencije je ugodno (*engl. Using Gen AI is enjoyable*).

S obzirom na to da se primarno tumačenje rezultata istraživanja vrši na temelju cjelovitog uzorka, unutar analize ovog poduzorka (UK) neće biti promjena u konceptualnom modelu istraživanja, no bitno je za napomenuti ovakava zapažanja u analizi multikolinearnosti među česticama. Sve ostale vrijednosti unutar mjernog modela upućuju prikladnost podataka, no može se uočiti kako su vrijednosti unutar konstrukta cijena također relativno visoke.

Tablica 27. Statistički pokazatelji i validacija mjernih varijabli [UK uzorak]

	Prosjek	Standardna devijacija	Vanjsko opterećenje	VIF
Namjera ponašanja				

BI1	5,157	1,431	0,877	2,564
BI2	4,963	1,478	0,862	2,132
BI3	5,532	1,285	0,940	3,626
Očekivani napori				
EE1	5,142	1,484	0,920	4,131
EE2	5,634	1,215	0,886	2,545
EE3	5,617	1,103	0,905	3,451
EE4	5,764	1,111	0,884	3,021
Olakšavajući uvjeti				
FC1	5,398	1,264	0,877	2,880
FC2	3,900	1,386	0,896	2,873
FC3	3,848	1,398	0,809	1,765
FC4	3,799	1,383	0,554	1,150
Hedonistička motivacija				
HM1	5,811	1,043	0,967	6,928
HM2	5,659	1,199	0,956	5,219
HM3	5,754	1,068	0,940	4,515
Navika				
HT1	4,988	1,465	0,901	2,813
HT2	5,266	1,379	0,833	2,405
HT3	5,376	1,318	0,851	2,498
HT4	5,197	1,364	0,847	2,116
Očekivani učinak				
PE1	4,475	1,318	0,857	2,119
PE2	4,629	1,355	0,889	2,802
PE3	4,597	1,354	0,867	2,493
PE4	3,632	1,898	0,876	2,571
Osobna inovativnost				
PI1	2,450	1,616	0,933	3,706
PI2	2,960	1,740	0,890	2,325
PI3	4,296	1,743	0,927	3,625
Cijena				
PV1	4,828	1,543	0,948	4,648
PV2	4,221	1,744	0,954	4,810
PV3	5,206	1,452	0,951	4,967
Društveni utjecaj				
SI1	5,818	1,051	0,930	3,442
SI2	5,866	1,035	0,926	3,513
SI3	5,796	0,992	0,945	3,943
Povjerenje				
TR1	5,739	1,081	0,922	4,247
TR2	5,672	1,226	0,884	3,264
TR3	4,182	1,777	0,860	2,112
TR4	5,211	1,504	0,882	2,730

Korištenje				
USE1	2,649	1,790	0,862	2,120
USE2	4,373	1,811	0,901	2,459
USE3	3,460	1,329	0,914	2,524

Cjeloviti naziv kodiranih čestica pojašnjen u Prilogu 1

VIF $\leq$ 5,0

Vanjsko opterećenje $\geq$ 0,6

Tablica 28., ispod teksta, prikazuje unakrsna faktorska opterećenja svih čestica u mjernom modelu. Kako je spomenuto i u analizi na cjelovitom uzorku, kod analize vanjskih opterećenja čestica uočeno je samo jedno odstupanje u čitavom modelu i riječ je o čestici FC4, ali s obzirom na to da ta čestica ne narušava stabilnost konstrukta, ali i radi bolje usporedivosti podataka s drugim istraživanjima, navedena je čestica ostavljena u mjernom modelu.

Tablica 28. Unakrsna faktorska opterećenja konstrukta [UK uzorak]

	BI	EE	FC	HM	HT	PE	PI	PV	SI	TR	USE
BI1	0,877	0,444	0,379	0,588	0,543	0,679	0,467	0,377	0,367	0,573	0,566
BI2	0,862	0,324	0,308	0,521	0,697	0,601	0,478	0,407	0,488	0,344	0,629
BI3	0,940	0,448	0,364	0,564	0,651	0,703	0,499	0,417	0,426	0,507	0,672
EE1	0,384	0,920	0,707	0,393	0,357	0,475	0,410	0,196	0,282	0,345	0,406
EE2	0,469	0,886	0,634	0,525	0,402	0,519	0,450	0,266	0,283	0,410	0,469
EE3	0,372	0,905	0,655	0,441	0,343	0,458	0,360	0,216	0,260	0,373	0,402
EE4	0,390	0,884	0,655	0,403	0,394	0,493	0,407	0,212	0,292	0,331	0,447
FC1	0,284	0,645	0,877	0,327	0,316	0,368	0,355	0,179	0,204	0,328	0,364
FC2	0,362	0,754	0,896	0,374	0,365	0,447	0,416	0,233	0,220	0,321	0,411
FC3	0,339	0,548	0,809	0,337	0,332	0,445	0,354	0,246	0,206	0,429	0,363
FC4	0,246	0,336	0,554	0,314	0,267	0,265	0,204	0,189	0,312	0,157	0,277
HM1	0,587	0,479	0,416	0,967	0,483	0,601	0,507	0,367	0,280	0,386	0,488
HM2	0,646	0,510	0,432	0,956	0,516	0,655	0,523	0,380	0,320	0,445	0,553
HM3	0,544	0,423	0,364	0,940	0,439	0,525	0,520	0,377	0,278	0,353	0,454
HT1	0,644	0,379	0,355	0,430	0,901	0,623	0,426	0,393	0,420	0,295	0,658
HT2	0,498	0,251	0,261	0,336	0,833	0,441	0,326	0,315	0,366	0,189	0,509
HT3	0,576	0,251	0,278	0,390	0,851	0,481	0,349	0,420	0,481	0,272	0,496
HT4	0,679	0,509	0,465	0,545	0,847	0,669	0,486	0,388	0,397	0,414	0,663
PE1	0,707	0,481	0,394	0,586	0,621	0,857	0,432	0,411	0,409	0,454	0,635
PE2	0,625	0,447	0,396	0,546	0,586	0,889	0,391	0,382	0,364	0,406	0,581
PE3	0,591	0,509	0,514	0,537	0,549	0,867	0,431	0,337	0,319	0,460	0,576
PE4	0,647	0,460	0,409	0,510	0,530	0,876	0,396	0,354	0,409	0,448	0,555
PI1	0,508	0,409	0,381	0,523	0,457	0,461	0,933	0,234	0,308	0,317	0,530
PI2	0,495	0,418	0,406	0,440	0,451	0,420	0,890	0,249	0,319	0,245	0,496
PI3	0,479	0,427	0,390	0,525	0,383	0,421	0,927	0,231	0,253	0,326	0,489
PV1	0,418	0,231	0,226	0,374	0,416	0,387	0,233	0,948	0,193	0,264	0,329

PV2	0,449	0,225	0,250	0,362	0,428	0,412	0,274	0,954	0,235	0,287	0,372
PV3	0,411	0,260	0,288	0,385	0,421	0,420	0,231	0,951	0,251	0,310	0,345
SI1	0,444	0,325	0,312	0,307	0,433	0,414	0,321	0,261	0,930	0,242	0,363
SI2	0,414	0,243	0,239	0,262	0,458	0,375	0,261	0,193	0,926	0,189	0,365
SI3	0,477	0,300	0,253	0,291	0,465	0,421	0,312	0,212	0,945	0,223	0,377
TR1	0,439	0,354	0,379	0,320	0,271	0,405	0,278	0,231	0,163	0,922	0,331
TR2	0,406	0,317	0,316	0,326	0,261	0,417	0,241	0,227	0,186	0,884	0,313
TR3	0,538	0,399	0,365	0,412	0,375	0,493	0,319	0,316	0,274	0,860	0,405
TR4	0,475	0,366	0,347	0,400	0,315	0,468	0,295	0,280	0,192	0,882	0,339
USE1	0,535	0,336	0,335	0,397	0,549	0,535	0,457	0,289	0,294	0,256	0,862
USE2	0,651	0,426	0,385	0,478	0,592	0,613	0,480	0,301	0,386	0,420	0,901
USE3	0,670	0,512	0,472	0,519	0,688	0,647	0,534	0,385	0,369	0,370	0,914

Pokazatelji unutarnje konzistentnosti i konvergentne valjanosti na razini konstrukta prikazane su u tablici 29. Vrijednost Cronbachove alfe i kompozitne pouzdanosti trebale bi biti 0,70 ili veće, dok je za vrijednosti AVE pokazatelja bitno da budu 0,50 ili veće. Unutar tablica je vidljivo kako niti jedna vrijednost nije narušena, a osobito je bitno za spomenuti kako konstrukt olakšavajućih uvjeta (FC) nema narušenu unutarnju konzistentnost niti konvergentnu valjanost ($\alpha = 0,882$; CR (rho_c) = 0,893; CR (rho_a) = 0,918; AVE = 0,737) usprkos uočenoj manjkavosti s vanjskim opterećenjem čestice FC4.

Tablica 29. Pokazatelji unutarnje konzistentnosti i konvergentne valjanosti [UK uzorak]

Konstrukt	Cronbach alfa (α)	CR (rho_c)	CR (rho_a)	AVE
Namjera ponašanja (BI)	0,873	0,877	0,922	0,799
Očekivani učinak (PE)	0,921	0,928	0,944	0,808
Očekivani naponi (EE)	0,794	0,823	0,870	0,634
Društveni utjecaj (SI)	0,951	0,958	0,968	0,911
Olakšavajući uvjeti (FC)	0,882	0,893	0,918	0,737
Hedonistička motivacija (HM)	0,896	0,898	0,927	0,761
Cijena (PV)	0,905	0,906	0,941	0,841
Navika (HT)	0,948	0,950	0,966	0,905
Osobna inovativnost (PI)	0,927	0,931	0,953	0,872
Povjerenje (TR)	0,910	0,917	0,937	0,788
Korištenje (USE)	0,873	0,885	0,922	0,797

AVE $\geq 0,5$

CR (rho_a; rho_c), Cronbach's $\geq 0,7$

Za procjenu diskriminantne valjanosti korišteni su Fornell-Larckerov kriterij i HTMT odnos. Rezultati analize pokazuju kako su svi korijeni pokazatelja AVE, prikazani na dijagonalni matrice, veći od odgovarajućih međukonstruktnih korelacija, čime je dodatno potvrđena jasnoća i međusobna razlika između promatranih konstrukta. Ponovno je vidljivo kako je sličnost između EE i FC relativno visoka, ali ne toliko da bi bila narušena diskriminantna valjanost.

Tablica 30. Diskriminantna valjanost prema Fornell-Larckerovom kriteriju [UK uzorak]

	BI	EE	FC	HM	HT	PE	PI	PV	SI	TR	USE
BI	0,894										
EE	0,454	0,899									
FC	0,392	0,737	0,796								
HM	0,624	0,495	0,425	0,954							
HT	0,707	0,419	0,406	0,504	0,858						
PE	0,740	0,543	0,488	0,626	0,657	0,872					
PI	0,539	0,456	0,428	0,541	0,470	0,473	0,917				
PV	0,448	0,250	0,268	0,392	0,444	0,427	0,259	0,951			
SI	0,478	0,311	0,288	0,308	0,484	0,433	0,320	0,238	0,934		
TR	0,530	0,409	0,398	0,416	0,350	0,507	0,323	0,302	0,234	0,887	
USE	0,698	0,483	0,450	0,525	0,687	0,674	0,551	0,367	0,395	0,396	0,893

HTMT kriterij kao suvremeniji i osjetljiviji pristup mjerenja diskriminantne valjanosti također pokazuje kako ne postoje pokazatelji narušene diskriminantne valjanost u mjernom modelu. Sve vrijednosti su ispod praga od 0,90 što potvrđuje dovoljnu razliku među promatranim konstruktima. Ponovno je vidljiva visoka sličnost između EE i FC, ali i PE i BI, no navedene vrijednosti i dalje su u granicama zadovoljavajuće razine. Ovakvi rezultati daju čvrstu osnovu za daljnju interpretaciju mjernog modela.

Tablica 31. Procjena diskriminantne valjanosti korištenjem HTMT kriterija [UK uzorak]

	BI	EE	FC	HM	HT	PE	PI	PV	SI	TR	USE
BI											
EE	0,501										
FC	0,471	0,850									
HM	0,682	0,521	0,493								
HT	0,794	0,447	0,475	0,538							
PE	0,833	0,595	0,578	0,672	0,722						
PI	0,606	0,496	0,499	0,583	0,516	0,525					
PV	0,492	0,265	0,311	0,414	0,483	0,461	0,279				
SI	0,530	0,335	0,349	0,326	0,536	0,471	0,348	0,254			
TR	0,588	0,438	0,459	0,439	0,374	0,555	0,352	0,320	0,249		
USE	0,793	0,527	0,535	0,568	0,766	0,756	0,617	0,399	0,435	0,433	

HTMT $\leq$ 0,90

7.4.2.1. Testiranje odnosa u modelu na uzorku iz Ujedinjenog Kraljevstva

Testiranje glavnih hipoteza ovog doktorskog rada kao i njegovog istraživačkog pitanja obrađena su na razini cjelokupnog uzorka u potpoglavlju 7.3.1., a analiza na poduzorcima testirat će se po istom principu kao i glavne hipoteze, radi usporedivosti, no neće se dodatno tumačiti kao hipoteze već samo kao odnosi unutar modela.

Iz tablice 32. može se uočiti kako je očekivani učinak (PE) najsnažniji prediktor namjere (BI) ($\beta = 0,316$; $t = 4,383$; $p < 0,05$) ponašanja, ali dob i spol kao moderatori ne utječu na odnos između PE i BI.

Nadalje, EE nije statistički značajan prediktor namjere ponašanja ($\beta = -0,041$; $t = 0,588$; $p > 0,05$), a niti moderirajući efekti spola, doba i iskustva nisu potvrđeni kao statistički značajni.

Konstrukt društvenog utjecaja (SI) ima statistički značajan utjecaj na namjeru ponašanja (BI) ($\beta = 0,134$; $t = 2,691$; $p < 0,05$), ali dob, spol i iskustvo kao moderatori navedenog odnosa nisu statistički značajni.

Utjecaj olakšavajućih uvjeta (FC) nema statistički značajan utjecaj na namjeru ponašanja (BI) ($\beta = -0,018$; $t = 0,292$; $p > 0,05$), ali ima na korištenje (USE) ($\beta = 0,102$; $t = 2,794$; $p < 0,05$). Osim navedenih odnosa, niti jedan od moderatora navedenih odnosa nije se pokazao kao statistički značajan.

Hedonistička motivacija (HM) također je statistički značajan prediktor namjere ponašanja (BI) ($\beta = 0,179$; $t = 2,914$; $p < 0,05$), no niti kod tog odnosa moderirajući efekt dobi, spola i iskustva nije potvrđen kao statistički značajan.

Utjecaj cijene, odnosno konstrukta vrijednosti cijene (PV) na namjeru ponašanja (BI) nije se potvrdio statistički značajnim ($\beta = 0,017$; $t = 0,377$; $p > 0,05$), a statistički značajni nisu ni moderatori navedenog odnosa, dob i spol.

Utjecaj navike (HT) statistički je značajan i drugi je najsnažniji prediktor namjere ponašanja (BI) ($\beta = 0,230$; $t = 4,300$; $p < 0,05$) i drugi najsnažniji prediktor korištenja (USE) ($\beta = 0,366$; $t = 6,623$; $p < 0,05$). Ipak, dob, spol i iskustvo kao moderatori navedenog odnosa nisu se pokazali statistički značajnima.

Osobna inovativnost (PI) kao konstrukt s kojim je proširen izvorni UTAUT2 pokazao se kao statistički značajan prediktor namjere ponašanja (BI) ($\beta = 0,112$; $t = 2,102$; $p < 0,05$). Premda nije osobito snažan prediktor, na poduzorku ispitanika iz Ujedinjenog Kraljevstva pokazao se snažnijim nego na cjelovitom uzorku. Ipak, moderatori dob, spol i iskustvo nemaju statistički značajan utjecaj na navedeni odnos.

Povjerenje (TR) kao još jedan konstrukt s kojim je proširen izvorni UTAUT2 također se pokazao statistički značajnim prediktorom namjere ponašanja (BI) ($\beta = 0,184$; $t = 3,938$; $p < 0,05$), no moderatori dob, spol i iskustvo nisu se pokazali statistički značajnima.

Utjecaj namjere ponašanja (BI) na korištenje (USE) također se pokazao kao umjereno snažan i statistički značajan prediktor ($\beta = 0,369$; $t = 7,319$; $p < 0,05$), no iskustvo se pokazalo značajnim moderatorom navedenog odnosa.

Tablica 32. Rezultati testiranja odnosa u konceptualnom modelu [UK uzorak]

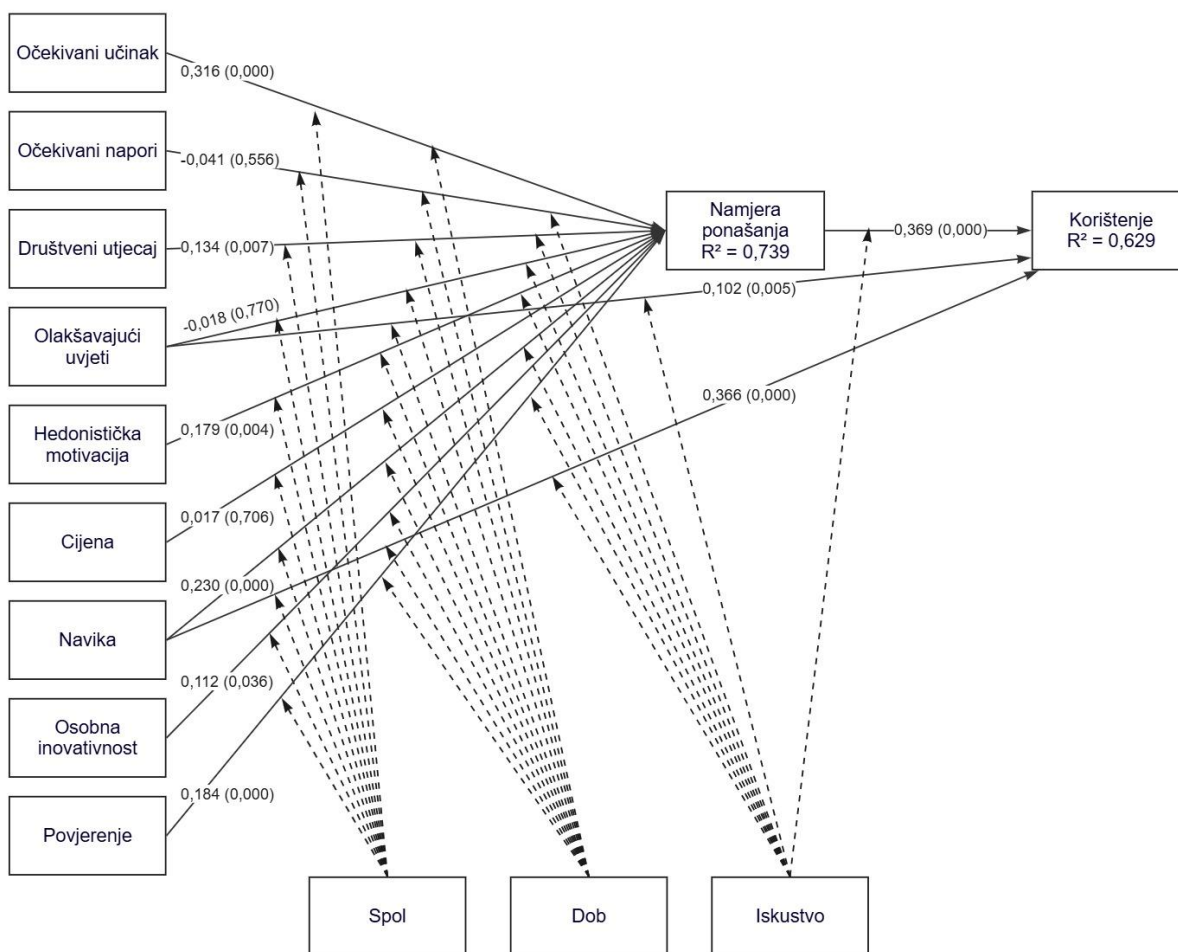
Hipoteza	Odnos	β	t-vrijednost	p	Ishod testiranja
H1	PE > BI	0,316	4,383	0,000***	Odnos je djelomično potvrđen.
	Dob x PE > BI	-0,042	0,937	0,349	
	Spol x PE > BI	-0,058	0,548	0,584	
H2	EE > BI	-0,041	0,588	0,556	Odnos nije potvrđen.
	Dob x EE > BI	0,002	0,043	0,966	
	Spol x EE > BI	0,135	1,308	0,191	
	Iskustvo x EE > BI	-0,028	0,643	0,520	
H3	SI > BI	0,134	2,691	0,007**	Odnos je djelomično potvrđen.
	Dob x SI > BI	-0,006	0,142	0,887	
	Spol x SI > BI	-0,051	0,680	0,497	
	Iskustvo x SI > BI	0,029	0,702	0,483	
H4a	FC > BI	-0,018	0,292	0,770	Odnos nije potvrđen.
	Dob x FC > BI	0,021	0,430	0,667	
	Spol x FC > BI	-0,130	1,440	0,150	
	Iskustvo x FC > BI	0,067	1,362	0,173	
H4b	FC > USE	0,102	2,794	0,005**	Odnos je djelomično potvrđen.
	Dob x FC > USE	-0,001	0,032	0,975	
	Iskustvo x FC > USE	0,015	0,461	0,645	
H5	HM > BI	0,179	2,914	0,004**	Odnos je djelomično potvrđen.
	Dob x HM > BI	-0,015	0,340	0,734	
	Spol x HM > BI	-0,101	1,019	0,308	
	Iskustvo x HM > BI	-0,044	0,982	0,326	
H6	PV > BI	0,017	0,377	0,706	Odnos nije potvrđen.
	Dob x PV > USE	-0,052	1,265	0,206	
	Spol x PV > USE	0,073	0,886	0,376	
H7a	HT > BI	0,230	4,300	0,000***	Odnos je djelomično potvrđen.
	Dob x HT > BI	-0,006	0,134	0,894	
	Spol x HT > BI	0,116	1,491	0,136	
	Iskustvo x HT > BI	-0,042	0,755	0,450	
H7b	HT > USE	0,366	6,623	0,000***	Odnos je djelomično potvrđen.
	Dob x HT > USE	0,056	1,352	0,177	
	Spol x HT > USE	-0,054	0,865	0,387	
	Iskustvo x HT > USE	-0,036	0,701	0,484	
H8	PI > BI	0,112	2,102	0,036*	Odnos je djelomično potvrđen.
	Dob x PI > BI	0,011	0,290	0,772	
	Spol x PI > BI	-0,028	0,365	0,715	
H9	TR > BI	0,184	3,938	0,000***	Odnos je djelomično potvrđen.
	Dob x TR > BI	0,026	0,674	0,500	
	Spol x TR > BI	-0,012	0,160	0,873	
	Iskustvo x TR > BI	0,003	0,092	0,926	

H10	BI > USE	0,369	7,319	0,000***	Odnos je djelomično potvrđen.
	Iskustvo x BI > USE	0,067	1,610	0,108	

β (Beta) = *Path coefficients* (standardizirani koeficijent putova); p = p vrijednost

* $p \leq 0,05$; ** $p < 0,01$; *** $p < 0,001$

Na slici 19. pojednostavljeno je prikazan odnos između nezavisnih i zavisnih varijabli navedenih u tablici 32., kao i vrijednosti standardiziranih koeficijenata putova i p-vrijednosti tih odnosa. Ipak, radi urednijeg prikaza nisu prikazane vrijednosti za moderatore, no u slučaju uzorka iz Ujedinjenog Kraljevstva niti jedan moderator nije statistički značajan. U konačnici, zanimljivo je za uočiti kako je objašnjenje varijance modela kod namjere ponašanja 73,9 % ($R^2 = 0,739$), dok je kod zavisne varijable korištenja objašnjenje varijance 62,9 % ($R^2 = 0,629$). S druge strane, Q^2 vrijednost za zavisnu varijablu BI iznosi 0,665, dok Q^2 za zavisnu varijablu iznosi 0,582. Obje vrijednosti su veće od 0,5 što ukazuje na izrazito snažnu prediktivnu sposobnost modela. Navedeni rezultati ukazuju na relativno veliku podudarnost između R^2 i Q^2 vrijednosti što ukazuje na stabilnu prediktivnu sposobnost modela u kojem nije korišteno pretjerano prilagođavanje podacima (engl. *overfitting*).



* koeficijent putanje - β (vrijednost ispred zagrade); p-vrijednost (vrijednost u zagradi)

Slika 19. Rezultati empirijskog istraživanja na konceptualnom modelu [UK uzorak]

Izvor: izrada autora

Ukoliko bi se promatrao model bez moderatora, SRMR vrijednost za cjeloviti model iznosi 0,057, dok za specificirani iznosi 0,059. Obje vrijednosti zadovoljavajuće su razine ($< 0,080$). Objašnjenje varijance kod namjere ponašanja je 70,6 % ($R^2 = 0,706$) dok je kod korištenja 58,2 % ($R^2 = 0,582$). Stone-Geisserov koeficijent predikcije za namjeru ponašanja iznosi 0,684, dok za korištenje iznosi 0,562. Ponovno je vidljivo kako su R^2 vrijednosti i Q^2 vrijednosti približno slične što upućuje na visoku stabilnost predikcije modela.

Dakle, u modelu bez moderatora objašnjenje varijance je za 3,3 postotna poena manja kod BI i 4,7 postotnih poena kod USE. Kako je vidljivo u tablici 33., ispod teksta, određeni odnosi promijenili su se u odnosu na model s uključenim moderatorima. Uočava se kako je bez moderirajućih efekata $FC > BI$ statistički značajan. Utjecaj očekivanih napora (EE) i cijene (PV) jedini nisu potvrđeni kao prediktori namjere ponašanja (BI) u modelu bez moderirajućih efekata.

Tablica 33. Testiranje odnosa u modelu bez moderatora [UK uzorak]

Odnos	β	t-vrijednost	p
PE > BI	0,291	6,286	0,000***
EE > BI	-0,000	0,009	0,993
SI > BI	0,102	2,622	0,009**
FC > BI	-0,096	2,035	0,042*
FC > USE	0,156	4,572	0,000***
HM > BI	0,142	3,027	0,002**
PV > BI	0,056	1,434	0,152
HT > BI	0,287	7,759	0,000***
HT > USE	0,348	6,993	0,000***
PI > BI	0,126	3,262	0,001**
TR > BI	0,180	5,042	0,000***
BI > USE	0,391	8,048	0,000***

β (Beta) = *Path coefficients* (standardizirani koeficijent putova); p = p vrijednost

* $p \leq 0,05$; ** $p < 0,01$; *** $p < 0,001$

7.5. Analiza rezultata s poduzorka iz Republike Hrvatske

7.5.1. Prikupljanje podataka i uzorak na tržištu Republike Hrvatske

U procesu regrutiranja ispitanika u Republici Hrvatskoj nije korištena platforma Prolific kao na tržištu Ujedinjenog Kraljevstva, već je autor samostalno prikupljao podatke jer platforma Prolific nema dovoljno veliku bazu ispitanika iz Hrvatske kako bi se moglo kvalitetno selektirati željenu ciljnu skupinu. Proces prikupljanja podataka trajao je od 5. ožujka do 27. ožujka 2025. godine. U procesu prikupljanja podataka kombinirano je nekoliko metoda među kojima je najznačajnija tehnika snježne grude (engl. *snowball*) ili tzv. lančano uzorkovanje, gdje je autor putem izravnih poruka e-pošte i direktnih poruka bliskim kontaktima prosljedio anketu uz zamolbu da prosljede svojim kontaktima. Svim primarnim kontaktima pružene su jasne upute o poželjnom profilu sudionika kako bi se ciljano usmjerilo širenje upitnika na relevantnu populaciju. Nadalje, u procesu prikupljanja ispitanika dodatnu su angažirani i studenti ispitivači koji su po jasno danim instrukcijama distribuirali istraživački upitnik ciljanom tržišnom segmentu. Ovakav pristup nalik je pristupu koji su koristili Gansser i Reich (2021), a koji jasno pojašnjava korištenje studenata kao produžene ruke istraživača u svrhu prikupljanja relevantnih ispitanika. Osim navedenih pristupa, autor je iskoristio i vlastitu privatnu i poslovnu bazu kontakata s ciljem dolaska do željene tržišne skupine. Dakle, može se zaključiti kako je osim tehnike snježne grude, korišteno i prigodno uzorkovanje te u određenoj mjeri i namjerno uzorkovanje uz posredovanje angažiranih ispitivača. Zbog korištenja i namjernog uzorka u prikupljanju podataka, podatak o postotku ispitanika koji koristi alate generativne umjetne inteligencije u odnosu na cjelokupni uzorak ne može biti interpretiran jer je cilj istraživanja ipak bio istražiti čimbenika prihvatanja i korištenja navedene tehnologije među korisnicima, odnosno onima koji su barem probali koristiti generativnu umjetnu inteligenciju.

Ukupno 1005 ispitanika iz Hrvatske pristupilo je upitniku, od čega 109 njih nije završilo ispunjavanje, 129 ispitanika je diskvalificirano na prvom pitanju jer su se izjasnili kako nikada nisu koristili alate umjetne inteligencije, a 43 ispitanika eliminirano je iz uzorka na temelju Alchemerove preporuke zbog upitne kvalitete prikupljenih odgovora, a koji sugerira isključivanje odgovora ispitanika ukoliko su u 1 % najbrže ispunjenih, imaju previše sličan obrazac ispunjavanja ili su nedosljedni u ispunjavanju. Od ukupno 1005 prikupljenih ispitanika, njih 767 prihvaćeno je sa zadovoljavajućom razinom kvalitete njihovih odgovora.

Kako bi rezultati istraživanja na tržištu Republike Hrvatske bili usporedivi s onima iz Ujedinjenog Kraljevstva, uzorak je oblikovan po prethodno navedenim kvotama te je reduciran na 400 ispitanika, slučajnim odabirom.

Tablica 34. Demografske karakteristike uzorka na tržištu Republike Hrvatske

Kategorije	Broj	Postotak (%)
SPOL		
Muškarci	218	54,5
Žene	182	45,5
DOBNE SKUPINE		
18 – 24	99	24,8
<i>Muškarci</i>	54	13,5
<i>Žene</i>	45	11,2
25 – 34	94	23,5
<i>Muškarci</i>	51	12,8
<i>Žene</i>	43	10,8
35 – 44	82	20,5
<i>Muškarci</i>	45	11,2
<i>Žene</i>	37	9,2
45 – 54	67	16,8
<i>Muškarci</i>	36	9,0
<i>Žene</i>	31	7,8
55 i više	58	14,5
<i>Muškarci</i>	32	8,0
<i>Žene</i>	26	6,5
RADNI STATUS		
Radim	259	64,8
Student sam	83	20,8
Ne radim i nisam student	12	3,0
Student sam, ali i radim	46	11,5
Idem u srednju školu	0	0,0
MJESTO STANOVANJA		
Ruralno područje	60	15,0

Mali grad (<i>do 15,000 stanovnika</i>)	44	11,0
Srednji grad (<i>do 50,000 stanovnika</i>)	68	17,0
Veliki grad (<i>više od 50,000 stanovnika</i>)	228	57,0
RAZINA OBRAZOVANJA		
Osnovna škola	0	0,0
Srednja škola	120	30,0
VSS / Stručni studiji	12	3,0
Prijediplomski studiji	84	21,0
Diplomski studiji	138	34,5
Poslijediplomski studiji	46	11,5

Među ispitanicima iz Hrvatske, najčešće korišteni alat generativne umjetne inteligencije je ChatGPT kojeg koristi čak 97 % korisnika navedene tehnologije. Nadalje, vrlo je popularan i Grammarly kojeg koristi 29,5 % ispitanika, Gemini koristi 26,5 % ispitanika, Perplexity koristi 22 % ispitanika. Slijede ih Microsoft Copilot (18,5 %), DeepSeek (18,5 %), Quillbot (9,5 %), Dall-E (8,8 %), Midjourney (7,8 %), Claude (6,5 %), Elicit (5,8 %), SciteAI (5,3 %), Grok (4,3 %), GitHub Copilot (4,3 %), Adobe Firefly (3,8 %), Google Notebook (3,8 %), Stable Diffusion (3 %), Meta Llama (2 %), Leonardo (2 %), Jasper AI (1,8 %), Mistral (1,3 %), HeyGen (1 %), Stability AI (0,8 %), AlphaCode (0,5 %), Runway (0,5 %), Poe (0,3 %). Kako je vidljivo iz uzorka Republike Hrvatske, značajno je veći postotak ispitanika koji imaju završen diplomski studij i poslijediplomski studij nego u uzorku iz Velike Britanije zbog čega se može uočiti razlika u alatima koje koriste, posebno se to može vidjeti kod onih alata koji pomažu u akademskim obvezama.

7.5.2. Rezultati primarnog istraživanja na tržištu Republike Hrvatske

Prvi korak u PLS-SEM analizi je provjera prikladnosti modela i tu se u pravilu koristi SRMR test čija bi vrijednost trebala biti ispod 0,08 kako bi se model smatrao prihvatljivim. U analizi podataka hrvatskih ispitanika SRMR vrijednost za cjeloviti model (*saturated*) iznosi 0,055, dok SRMR vrijednost kod specificiranog modela (*estimated*) iznosi 0,056. Oba rezultata upućuju na dobru pouzdanost modela.

Tablica 35., ispod teksta, prikazuje prosječne vrijednosti i standardnu devijaciju za svaku česticu mjernog modela, ali i vanjska opterećenja i faktor inflacije varijance. Iz priloženog je vidljivo kako je standardna devijacija kod svih čestica između 1,0 i 2,0 što ukazuje na umjerenu

do visoku raspršenost odgovora, što pak upućuje na dobru raznolikost među ispitanicima u uzorku. Kod stupca vanjskog opterećenja čestice ponovno se može uočiti kako je jedina problematična čestica FC4 koja iznosi 0,544, a što je ispod preporučenog praga od 0,6. Navedena čestica prethodno je protumačena na razini cjelovitog uzorka. Ipak, u odnosu na uzorak iz Ujedinjenog Kraljevstva, ovdje nema povišenih VIF pokazatelja kod konstrukta hedonističke motivacije, a niti kod drugih čestica iz modela.

Tablica 35. Statistički pokazatelji i validacija mjernih varijabli [RH uzorak]

	Prosjek	Standardna devijacija	Vanjsko opterećenje	VIF
Namjera ponašanja				
BI1	5,915	1,167	0,836	1,953
BI2	4,577	1,689	0,869	2,414
BI3	5,173	1,447	0,938	3,380
Očekivani napori				
EE1	5,787	1,236	0,903	3,087
EE2	5,662	1,161	0,836	2,059
EE3	5,795	1,130	0,862	2,331
EE4	5,537	1,176	0,889	2,834
Olakšavajući uvjeti				
FC1	5,670	1,312	0,856	1,960
FC2	5,492	1,302	0,826	1,798
FC3	5,425	1,241	0,820	1,570
FC4	5,082	1,446	0,544	1,210
Hedonistička motivacija				
HM1	5,803	1,224	0,929	3,530
HM2	5,702	1,206	0,895	2,438
HM3	5,945	1,154	0,932	3,604
Navika				
HT1	4,815	1,726	0,866	2,467
HT2	2,815	1,782	0,837	2,550
HT3	2,732	1,785	0,770	2,262
HT4	4,233	1,797	0,892	2,710
Očekivani učinak				
PE1	5,888	1,221	0,866	2,415
PE2	5,445	1,344	0,878	2,497
PE3	5,915	1,311	0,888	2,807
PE4	5,442	1,453	0,863	2,255
Osobna inovativnost				
PI1	4,348	1,821	0,928	3,146
PI2	4,045	1,882	0,875	2,214
PI3	5,015	1,688	0,916	2,963

Cijena				
PV1	5,025	1,321	0,928	3,388
PV2	4,787	1,309	0,884	2,359
PV3	5,043	1,317	0,932	3,459
Društveni utjecaj				
SI1	4,338	1,519	0,918	2,952
SI2	4,242	1,558	0,933	3,574
SI3	4,345	1,512	0,929	3,472
Povjerenje				
TR1	5,955	1,176	0,898	3,349
TR2	6,003	1,132	0,913	3,629
TR3	5,883	1,195	0,903	3,068
TR4	5,840	1,181	0,887	2,875
Korištenje				
USE1	2,390	1,685	0,790	1,582
USE2	4,530	1,743	0,887	2,129
USE3	3,705	1,419	0,896	2,146

Cjeloviti naziv kodiranih čestica pojašnjen u Prilogu 1

VIF $\leq$ 5,0

Vanjsko opterećenje $\geq$ 0,6

Tablica 36., koja prikazuje sva faktorska opterećenja čestica u radu, prikazuje kako je čestica FC4 ponovno jedina problematična jer ne prelazi preporučeni prag od 0,6. Ipak vidljivo je i kako postoji velika sličnost između čestica u konstrukt EE i konstrukt FC. Navedena su opažanja uočena i na razini cjelokupnog uzorka. Ipak, ona ne narušavaju valjanost mjernog modela u mjeri da je potrebno raditi modifikacije unutar konceptualnog modela.

Tablica 36. Unakrsna faktorska opterećenja konstrukta [RH uzorak]

	BI	EE	FC	HM	HT	PE	PI	PV	SI	TR	USE
BI1	0,836	0,394	0,476	0,593	0,485	0,553	0,418	0,410	0,331	0,615	0,471
BI2	0,869	0,312	0,294	0,516	0,649	0,474	0,474	0,432	0,414	0,373	0,540
BI3	0,938	0,358	0,409	0,559	0,683	0,526	0,513	0,477	0,397	0,500	0,610
EE1	0,363	0,903	0,632	0,473	0,318	0,434	0,354	0,339	0,210	0,319	0,282
EE2	0,336	0,836	0,558	0,519	0,330	0,457	0,278	0,399	0,319	0,265	0,325
EE3	0,339	0,862	0,568	0,463	0,259	0,368	0,186	0,338	0,258	0,278	0,217
EE4	0,361	0,889	0,641	0,471	0,285	0,402	0,317	0,319	0,288	0,307	0,257
FC1	0,373	0,560	0,856	0,378	0,246	0,342	0,272	0,310	0,192	0,374	0,298
FC2	0,327	0,666	0,826	0,364	0,281	0,289	0,271	0,315	0,135	0,268	0,318
FC3	0,425	0,523	0,820	0,407	0,345	0,379	0,339	0,400	0,295	0,404	0,308
FC4	0,196	0,336	0,544	0,297	0,217	0,178	0,097	0,261	0,222	0,172	0,107
HM1	0,578	0,518	0,423	0,929	0,403	0,510	0,342	0,423	0,310	0,492	0,355
HM2	0,573	0,537	0,460	0,895	0,381	0,498	0,352	0,445	0,334	0,482	0,365
HM3	0,582	0,463	0,408	0,932	0,402	0,534	0,359	0,414	0,344	0,521	0,363
HT1	0,701	0,363	0,384	0,498	0,866	0,528	0,466	0,496	0,396	0,323	0,664
HT2	0,479	0,187	0,214	0,236	0,837	0,317	0,376	0,323	0,296	0,170	0,517
HT3	0,427	0,196	0,156	0,202	0,770	0,268	0,319	0,322	0,336	0,152	0,422
HT4	0,653	0,353	0,378	0,433	0,892	0,494	0,532	0,397	0,388	0,347	0,647
PE1	0,498	0,405	0,357	0,491	0,408	0,866	0,339	0,362	0,358	0,442	0,412
PE2	0,529	0,425	0,344	0,480	0,484	0,878	0,337	0,454	0,416	0,383	0,472
PE3	0,485	0,435	0,371	0,488	0,435	0,888	0,312	0,391	0,344	0,414	0,435
PE4	0,532	0,397	0,326	0,497	0,410	0,863	0,383	0,422	0,395	0,414	0,426
PI1	0,513	0,288	0,292	0,344	0,476	0,353	0,928	0,322	0,219	0,363	0,435
PI2	0,448	0,275	0,269	0,300	0,470	0,373	0,875	0,240	0,276	0,341	0,505
PI3	0,484	0,326	0,359	0,393	0,456	0,346	0,916	0,269	0,203	0,428	0,460
PV1	0,464	0,346	0,397	0,409	0,412	0,414	0,287	0,928	0,302	0,368	0,338
PV2	0,426	0,364	0,349	0,412	0,411	0,416	0,282	0,884	0,285	0,314	0,338
PV3	0,478	0,385	0,401	0,455	0,457	0,453	0,276	0,932	0,308	0,397	0,367

SI1	0,406	0,310	0,270	0,330	0,398	0,412	0,255	0,273	0,918	0,269	0,326
SI2	0,401	0,243	0,241	0,320	0,372	0,376	0,225	0,291	0,933	0,267	0,302
SI3	0,393	0,300	0,232	0,347	0,410	0,420	0,228	0,344	0,929	0,238	0,356
TR1	0,476	0,281	0,368	0,468	0,269	0,418	0,359	0,340	0,248	0,898	0,292
TR2	0,514	0,314	0,395	0,534	0,279	0,464	0,413	0,348	0,249	0,913	0,311
TR3	0,540	0,313	0,368	0,474	0,307	0,433	0,405	0,405	0,275	0,903	0,345
TR4	0,482	0,298	0,354	0,478	0,262	0,385	0,318	0,324	0,228	0,887	0,299
USE1	0,406	0,118	0,207	0,266	0,495	0,323	0,308	0,274	0,252	0,244	0,790
USE2	0,594	0,275	0,308	0,381	0,595	0,463	0,466	0,334	0,342	0,357	0,887
USE3	0,566	0,372	0,376	0,355	0,659	0,482	0,520	0,363	0,310	0,287	0,896

Tablica 37. prikazuje pokazatelje unutarnje konzistentnosti i konvergentne valjanosti. Vrijednosti Cronbachove alfe i kompozitne pouzdanosti trebale bi biti 0,70 ili veće i to je slučaj kod svih konstrukta. S druge strane, vrijednost pokazatelja AVE trebala bi biti 0,50 ili veća i tu je također slučaj da su sve vrijednosti zadovoljavajuće razine. Može se uočiti kako su vrijednosti kod konstrukta očekivanih napora (EE) najslabije ($\alpha = 0,772$; CR (rho_c) = 0,821; CR (rho_a) = 0,852; AVE = 0,596), no navedene vrijednosti nisu zabrinjavajuće jer su i dalje u granicama zadovoljavajuće razine.

Tablica 37. Pokazatelji unutarnje konzistentnosti i konvergentne valjanosti [RH uzorak]

Konstrukt	Cronbach alfa (α)	CR (rho_c)	CR (rho_a)	AVE
Namjera ponašanja	0,856	0,865	0,913	0,778
Očekivani učinak	0,896	0,897	0,928	0,762
Očekivani naponi	0,772	0,821	0,852	0,596
Društveni utjecaj	0,907	0,907	0,942	0,844
Olakšavajući uvjeti	0,866	0,893	0,907	0,710
Hedonistička motivacija	0,897	0,898	0,928	0,764
Cijena	0,891	0,896	0,932	0,822
Navika	0,903	0,907	0,939	0,838
Osobna inovativnost	0,918	0,918	0,948	0,859
Povjerenje	0,922	0,925	0,945	0,811
Korištenje	0,822	0,841	0,894	0,738

AVE $\geq$ 0,5

CR (rho_a; rho_c), Cronbach's $\geq$ 0,7

Za procjenu diskriminantne valjanosti korišteni su Fornell-Larckerov kriterij i HTMT kriterij. Rezultati analize pokazuju kako su korijeni pokazatelja AVE, prikazani na dijagonali matrice, veći od odgovarajućih međukonstruktnih korelacija, čime je dodatno potvrđena jasnoća i međusobna razlika između promatranih konstrukta. Ponovno je vidljiva sličnost između konstrukta EE i FC, ali navedena sličnost nije dovoljno velika da bi narušila diskriminantnu valjanost.

Tablica 38. Diskriminantna valjanost prema Fornell-Larckerovom kriteriju [RH uzorak]

	BI	EE	FC	HM	HT	PE	PI	PV	SI	TR	USE
BI	0,882										
EE	0,401	0,873									
FC	0,444	0,688	0,772								
HM	0,629	0,551	0,468	0,919							
HT	0,691	0,341	0,355	0,431	0,843						
PE	0,585	0,475	0,399	0,560	0,498	0,874					
PI	0,532	0,327	0,339	0,382	0,515	0,393	0,906				
PV	0,499	0,399	0,419	0,465	0,467	0,468	0,307	0,915			
SI	0,432	0,307	0,267	0,358	0,424	0,434	0,255	0,326	0,927		
TR	0,560	0,335	0,413	0,543	0,311	0,473	0,417	0,395	0,279	0,900	
USE	0,615	0,309	0,354	0,393	0,684	0,500	0,512	0,380	0,353	0,347	0,859

HTMT kriterij kao suvremeniji i osjetljiviji pristup također ukazuje na činjenicu da ne postoji narušena diskriminantna valjanost među konstruktima. Sve promatrane vrijednosti su ispod preporučenog praga od 0,90 što potvrđuje dovoljno veliku razliku među promatranim konstruktima. Ponovno se može uočiti relativno visoka sličnost između konstrukta EE i konstrukta FC, no i dalje je u granicama zadovoljavajuće razine. Navedeni pokazatelji daju čvrstu osnovu za daljnju interpretaciju mjernog modela.

Tablica 39. Procjena diskriminantne valjanosti korištenjem HTMT kriterija [RH uzorak]

	BI	EE	FC	HM	HT	PE	PI	PV	SI	TR	USE
BI											
EE	0,459										
FC	0,529	0,813									
HM	0,716	0,612	0,561								
HT	0,773	0,370	0,409	0,457							
PE	0,669	0,531	0,464	0,620	0,539						
PI	0,607	0,364	0,382	0,424	0,571	0,440					
PV	0,566	0,445	0,499	0,514	0,514	0,518	0,341				
SI	0,487	0,340	0,325	0,393	0,470	0,477	0,284	0,359			
TR	0,632	0,368	0,469	0,593	0,328	0,519	0,457	0,430	0,302		
USE	0,722	0,346	0,411	0,450	0,783	0,573	0,590	0,437	0,404	0,395	

HTMT $\leq$ 0,90

7.5.2.1. Testiranje odnosa u modelu na uzorku iz Republike Hrvatske

Kako je prethodno spomenuto, testiranje glavnih hipoteza ovog doktorskog rada kao i njegovog istraživačkog pitanja obrađeno je na razini cjelokupnog uzorka u potpoglavlju 7.3., a analiza na poduzorcima testirat će se po istom principu kao glavne hipoteze, radi usporedivosti, no neće se dodatno tumačiti kao hipoteze već samo kao odnosi unutar modela.

Iz tablice 40., ispod teksta, može se vidjeti kako očekivani učinak (PE) nije statistički značajan prediktor namjere ponašanja (BI) ($\beta = -0,030$; $t = 0,634$; $p > 0,05$). Za razliku od uzorka iz UK-

a gdje je PE bio najsnažniji prediktor namjere ponašanje, ovdje je vrijednost standardiziranog koeficijenta puta negativna i nije statistička značajna, no prediktori dob ($\beta = 0,100$; $t = 2,399$; $p < 0,05$) i spol ($\beta = 0,265$; $t = 2,810$; $p < 0,05$) jesu statistički moderatori navedenog odnosa, pri čemu je utjecaj očekivanog učinka na namjeru ponašanja izraženiji kod mladih i žena.

Nastavno, kao i kod cjelokupnog uzorka, ali i UK uzorka, niti ovdje se konstrukt očekivanih napora (EE) nije potvrdio kao statistički značajan prediktor namjere ponašanja (BI) ($\beta = 0,028$; $t = 0,461$; $p > 0,05$). S druge strane, kao i kod odnosa PE > BI, ovdje su se također dva moderatora potvrdila statistički značajnima, a to su spol ($\beta = -0,211$; $t = 2,246$; $p < 0,05$) i iskustvo ($\beta = 0,131$; $t = 2,703$; $p < 0,05$), ali dob nije. Utjecaj očekivanih napora (EE) snažniji je kod žena i kod iskusnijih ispitanika što je suprotno početnim pretpostavkama.

Društveni utjecaj (SI) također se nije potvrdio kao prediktor namjere ponašanja (BI) ($\beta = 0,043$; $t = 0,799$; $p > 0,05$). Statistički značajnim nisu se potvrdili niti moderatori dob, spol i iskustvo.

Statistički značajnim prediktorom namjere ponašanja nije se potvrdio ni konstrukt olakšavajućih uvjeta (FC) ($\beta = 0,017$; $t = 0,273$; $p > 0,05$). Također, ni moderatori dob, spol i iskustvo nisu se pokazali statistički značajnima. Ipak, za razliku od rezultata na cjelovitom uzorku, ovdje se olakšavajući uvjeti (FC) nisu potvrdili kao značajan prediktor korištenja (USE) ($\beta = 0,071$; $t = 1,705$; $p > 0,05$). Dob i spol također nisu potvrđeni kao značajni moderatori navedenog odnosa.

Nadalje, hedonistička motivacija (HM) također se pokazala statistički značajnom u utjecaju na namjeru ponašanja (BI) ($\beta = 0,309$; $t = 4,561$; $p < 0,05$). Ipak, dob, spol i iskustvo nisu potvrđeni kao statistički značajni moderatori spomenutog odnosa.

Cijena, kao i u analizi cjelovitog uzorka i UK uzorka, nije potvrđena kao prediktor namjere ponašanja (BI) ($\beta = 0,044$; $t = 0,887$; $p > 0,05$). Značajni nisu niti moderatori dob i spol.

Navika (HT) se u slučaju RH uzorka potvrdila kao statistički značajan, ali i najsnažniji prediktor namjere ponašanja (BI) ($\beta = 0,387$; $t = 6,573$; $p < 0,05$) i korištenja ($\beta = 0,474$; $t = 7,815$; $p < 0,05$). Ipak, niti jedan moderator u promatranim odnosima nije se pokazao statistički značajnim.

U slučaju RH uzorka, vidljivo je kako osobna inovativnost (PI) nije statistički značajan prediktor namjere ponašanja ($\beta = 0,047$; $t = 0,809$; $p > 0,05$) premda je kod UK uzorka, ali i na razini cjelokupnog uzorka bio statistički značajan. Kod promatranog odnosa, statistički značajni nisu niti moderatori dob i spol.

Povjerenje (TR) ima statistički značajan utjecaj na namjeru ponašanja (BI) ($\beta = 0,170$; $t = 3,183$; $p < 0,05$), ali dob, spol i iskustvo nisu potvrđeni kao značajni moderatori spomenutog odnosa.

U konačnici, podatci pokazuju kako namjera ponašanja (BI) ima statistički značajan utjecaj na stvarno korištenje (USE) ($\beta = 0,233$; $t = 4,726$; $p < 0,05$), ali da iskustvo nije značajan moderator navedenog odnosa.

Tablica 40. Rezultati testiranja odnosa u konceptualnom modelu [RH uzorak]

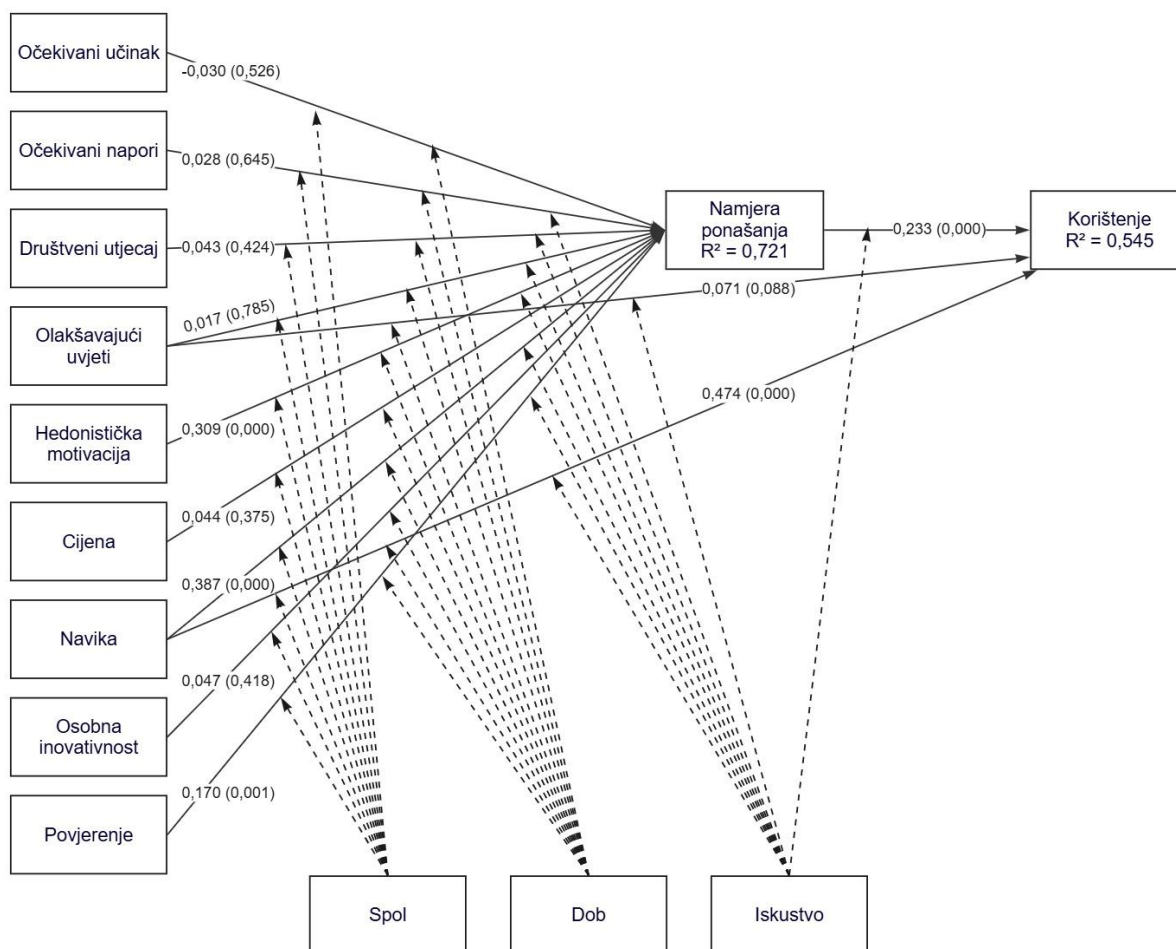
Hipoteza	Odnos	β	t-vrijednost	p	Ishod testiranja
H1	PE > BI	-0,030	0,634	0,526	Odnos nije potvrđen.
	Dob x PE > BI	0,100	2,399	0,016*	
	Spol x PE > BI	0,265	2,810	0,005**	
H2	EE > BI	0,028	0,461	0,645	Odnos nije potvrđen.
	Dob x EE > BI	-0,011	0,229	0,819	
	Spol x EE > BI	-0,211	2,246	0,025*	
	Iskustvo x EE > BI	0,131	2,703	0,007**	
H3	SI > BI	0,043	0,799	0,424	Odnos nije potvrđen.
	Dob x SI > BI	-0,027	0,668	0,504	
	Spol x SI > BI	-0,027	0,362	0,718	
	Iskustvo x SI > BI	0,013	0,309	0,757	
H4a	FC > BI	0,017	0,273	0,785	Odnos nije potvrđen.
	Dob x FC > BI	0,021	0,406	0,685	
	Spol x FC > BI	0,078	0,840	0,401	
	Iskustvo x FC > BI	-0,066	1,378	0,168	
H4b	FC > USE	0,071	1,705	0,088	Odnos nije potvrđen.
	Dob x FC > USE	0,005	0,142	0,887	
	Iskustvo x FC > USE	0,051	1,357	0,175	
H5	HM > BI	0,309	4,561	0,000***	Odnos je djelomično potvrđen.
	Dob x HM > BI	-0,026	0,508	0,611	
	Spol x HM > BI	-0,150	1,386	0,166	
	Iskustvo x HM > BI	-0,026	0,480	0,631	
H6	PV > BI	0,044	0,887	0,375	Odnos nije potvrđen.
	Dob x PV > USE	-0,004	0,104	0,917	
	Spol x PV > USE	0,033	0,434	0,664	
H7a	HT > BI	0,387	6,573	0,000***	Odnos je djelomično potvrđen.
	Dob x HT > BI	-0,006	0,152	0,879	
	Spol x HT > BI	0,027	0,325	0,745	
	Iskustvo x HT > BI	-0,024	0,675	0,499	
H7b	HT > USE	0,474	7,815	0,000***	Odnos je djelomično potvrđen.
	Dob x HT > USE	0,003	0,082	0,935	
	Spol x HT > USE	-0,036	0,527	0,598	
	Iskustvo x HT > USE	0,059	1,409	0,159	
H8	PI > BI	0,047	0,809	0,418	Odnos nije potvrđen.
	Dob x PI > BI	-0,049	1,243	0,214	
	Spol x PI > BI	0,123	1,603	0,109	
H9	TR > BI	0,170	3,183	0,001**	

	Dob x TR > BI	-0,005	0,105	0,916	Odnos je djelomično potvrđen.
	Spol x TR > BI	-0,020	0,221	0,825	
	Iskustvo x TR > BI	-0,048	1,108	0,268	
H10	BI > USE	0,233	4,726	0,000***	Odnos je djelomično potvrđen.
	Iskustvo x BI > USE	0,004	0,076	0,940	

β (Beta) = *Path coefficients* (standardizirani koeficijent putova); p = p vrijednost

* $p \leq 0,05$; ** $p < 0,01$; *** $p < 0,001$

Slika 20. prikazuje pojednostavljeni prikaz odnosa u modelu. Vrijednosti utjecaja moderatora na odnose između nezavisnih i zavisnih varijabli izostavljene su zbog urednijeg prikaza grafikona, no vrijednosti su vidljive u tablici 40. PLS-SEM analiza modela na promatranom uzorku ispitanika iz Republike Hrvatske potvrdila je objašnjenje varijance modela kod namjere ponašanja (BI) od 72,1 % ($R^2 = 0,721$) i 54,5 % ($R^2 = 0,545$) objašnjenja varijance kod zavisne varijable korištenja (USE). S druge strane, Stone-Geisserov koeficijent predikcije (Q^2) iznosi 0,645 za BI i 0,500 za USE. S obzirom na to da su obje vrijednosti veće ili jednake 0,50, podatci ukazuju na izrazito visoku prediktivnu relevantnost modela. Nadalje, vidljivo je i navedenih rezultata kako su vrijednosti između R^2 i Q^2 relativno približne, što ukazuje na visoku stabilnost predikcije u kojoj nije prisutan problem pretjeranog prilagođavanja podacima (engl. *overfitting*).



* koeficijent putanje - β (vrijednost ispred zagrade); p-vrijednost (vrijednost u zagradi)

Slika 20. Rezultati empirijskog istraživanja na konceptualnom modelu [RH uzorak]

Izvor: izrada autora

Ukoliko bi se promatrao model bez moderatora, SRMR vrijednost za cjeloviti model iznosi 0,057, dok za specificirani iznosi 0,058. Obje vrijednosti zadovoljavajuće su razine ($< 0,08$). Objašnjenje varijance kod namjere ponašanja je 67,8 % ($R^2 = 0,678$) dok je kod korištenja 51,2 % ($R^2 = 0,512$). Stone-Geisserov koeficijent predikcije za namjeru ponašanja iznosi 0,659, dok za korištenje iznosi 0,496. Ponovno je vidljivo kako su R^2 vrijednosti i Q^2 vrijednosti približno slične što upućuje na visoku stabilnost predikcije modela.

Dakle, u modelu bez moderatora objašnjenje varijance niža je za 4,3 postotna poena kod BI i 3,3 postotnih poena kod USE. Kako je vidljivo u tablici 41., ispod teksta, vidljive su promjene u odnosu između nezavisnih i zavisnih varijabli u odnosu na model s moderirajućim efektima. Vidljivo je kako je $PE > BI$ sada statistički značajan prediktor namjere ponašanja, jednako kao i $PI > BI$, što ukazuje da moderator oslabljuju prediktivnu snagu navedenih konstrukta.

Tablica 41. Testiranje odnosa u modelu bez moderatora (RH uzorak)

Odnos	β	t-vrijednost	p
PE > BI	0,099	2,365	0,018*
EE > BI	-0,092	1,851	0,064
SI > BI	0,054	1,450	0,147
FC > BI	0,072	1,555	0,120
FC > USE	0,071	1,781	0,075
HM > BI	0,243	4,682	0,000***
PV > BI	0,042	1,080	0,280
HT > BI	0,390	10,092	0,000***
HT > USE	0,489	10,548	0,000***
PI > BI	0,101	2,655	0,008**
TR > BI	0,187	4,521	0,000***
BI > USE	0,247	5,131	0,000***

β (Beta) = *Path coefficients* (standardizirani koeficijent putova); p = p vrijednost

* $p \leq 0,05$; ** $p < 0,01$; *** $p < 0,001$

8. RASPRAVA

Ovo poglavlje ima za cilj detaljno analizirati i interpretirati rezultate empirijskog istraživanja predstavljenog u prethodnom poglavlju, stavljajući ih u kontekst postojećih spoznaja o prihvatanju i korištenju tehnologija, s naglaskom na generativnu umjetnu inteligenciju. Dobiveni rezultati usporedit će se s nalazima drugih istraživanja na ovu i povezane teme. Također, unutar ovog poglavlja naglasak će biti na implikacijama i ograničenjima ovog istraživanja, ali i na preporukama za buduća istraživanja.

8.1. Metodološka razmatranja i evaluacija mjernog modela

Podatci za empirijsko istraživanje prikupljeni su na dva različita tržišta (Ujedinjeno Kraljevstvo i Republika Hrvatska) koristeći različite metode regrutacije ispitanika. Na tržištu Ujedinjenog Kraljevstva korištena je platforma Prolific, dok je na hrvatskom tržištu bila kombinacija nekoliko metoda među kojima je najznačajnija bila tehnika snježne grude (engl. *snowball*) i korištenje studenata kao ispitivača relevantnih segmentnih skupina po uzoru na Gansser i Reich (2021). Naglasak je stavljen na korisnike starije od 18 godina koji su barem pokušali koristiti alate generativne umjetne inteligencije. Platforma Prolific omogućava filtriranje željenih demografskih skupina, pri čemu je napravljen kvotni uzorak ispitanika koji su imali prethodna iskustva s nekim od alata generativne umjetne inteligencije. Hrvatski uzorak, gdje je korištena kombinacija metoda, može imati određenih nedostataka kao što je veća vjerojatnost sudjelovanja pojedinaca iz sličnih društvenih krugova. Ipak, primjenom kvotnog uzorka po demografskim karakteristikama dobi i spola te filtriranja na temelju kvalitete podataka (oslanjajući se na funkcionalnosti softvera Alchemer), autor je imao za cilj osigurati što kvalitetniji i usporediviji uzorak između spomenutih tržišta. Veličina uzorka od 400 ispitanika po zemlji smatra se prikladnom za PLS-SEM analizu, čak i za modele ovakve kompleksnosti (Hair i sur., 2019b).

Konceptualni model korišten u ovom doktorskom radu temelji se na proširenom UTAUT2 modelu koji uključuje sedam temeljnih konstrukta: očekivani učinak (PE), očekivani naponi (EE), društveni utjecaj (SI), olakšavajući uvjeti (FC), hedonistička motivacija (HM), vrijednost cijene (PV), navika (HT), koji je dodatno proširen s konstruktima osobne inovativnosti (PI) i povjerenje (TR). U radu se analizira utjecaj navedenih prediktora na namjeru ponašanja (BI) i stvarno korištenje (USE) uz moderirajuće učinke dobi, spola i iskustva.

Kako je navedeno u poglavlju 7.1., za analizu podataka u ovom doktorskom radu korištena je metoda parcijalnih najmanjih kvadrata (PLS-SEM). Odluka o primjeni PLS-SEM-a, umjesto standardnog modeliranja strukturalnim jednadžbama temeljenog na kovarijanci (CB-SEM), donesena je na temelju nekoliko ključnih karakteristika ovog istraživanja. Prije svega, navedeni konceptualni model ovog doktorskog rada uključuje značajan broj prediktora i moderatora, a u čemu je PLS-SEM efikasniji od CB-SEM-a. Drugi važan razlog je istraživanje predviđanja i objašnjavanje varijance zavisnih konstrukta, a ne samo potvrda teorije kao bitan cilj doktorskog rada (Hair i sur., 2017a). Pregled literature (poglavlje 4.) potvrđuje kako čak 71 % analiziranih studija koji koriste UTAUT2 primjenjuje PLS-SEM što zasigurno potvrđuje njegovu široku prihvaćenost među istraživačima u ovom znanstvenom području.

Nadalje, PLS-SEM ne postavlja stroge zahtjeve o normalnosti distribucije podataka, što je često prednost kod rada s anketnim podacima koji rijetko udovoljavaju strogim kriterijima normalnosti distribucije podataka. Iako je u ovom istraživanju korišten relativno velik uzorak ($N = 802$; $N_{UK} = 402$, $N_{RH} = 400$), PLS-SEM je prigodan za rad i s manjim uzorcima. Nadalje, PLS-SEM za testiranje značajnosti koeficijenta putanje koristi *bootstrapping* proceduru s 5000 uzoraka (Hair i sur., 2022), a što je primijenjeno i u ovom doktorskom radu. *Bootstrapping* procedura osigurava pouzdanu procjenu statističke značajnosti, čak i bez pretpostavke o normalnoj distribuciji podataka.

Iako koncept *model fit* pokazatelja, koji je standard za CB-SEM analize, nije preporučljiv za PLS-SEM analize (Hair i sur., 2021), određeni pokazatelji ipak mogu pružiti uvid u to koliko dobro model objašnjava empirijske podatke. U ovom istraživanju korišten je standardizirani korijen srednje kvadratne rezidualne pogreške (SRMR). Za cjelokupni model, SRMR vrijednost od 0,050 (*saturated*) i 0,052 (*estimated*) značajno je ispod preporučenog praga od 0,08 (Hu i Bentler, 1999; Henseler i sur., 2016), što ukazuje na dobru prilagođenost modela. Slične vrijednosti su i kod poduzoraka iz Hrvatske (0,054 i 0,056) i UK-a (0,055 i 0,057) što potvrđuje kako model prikladno opaža korelacije s podacima na oba tržišta. Važan je to metodološki nalaz koji sugerira da struktura odnosa postavljena modelom nije u raskoraku s empirijski utvrđenom strukturom.

Također, u radu je prikazana i evaluacija mjernih instrumenata koja osigurava kvalitetu istraživanja, a pri tome je naglasak stavljen na unutarnju konzistentnost, konvergentnu valjanost i diskriminantnu valjanost konstrukta. Unutarnja konzistentnost mjerena je pomoću Cronbachove alfe (α), Jöreskogove kompozitne pouzdanosti (CR rho_c) i Dijkstrine kompozitne pouzdanosti (CR rho_a). Kako je prikazano u tablici 21., unutar poglavlja 7.3., ali

i unutar odgovarajućih tablica kod poduzoraka u poglavljima 7.4.2. i 7.5.2., vrijednosti za sva tri pokazatelja značajno premašuju preporučeni prag od 0,70 (Hair i sur., 2021; Goforth, 2015). Dobiveni rezultati ukazuju na visoku razinu unutarnje konzistentnosti što znači da čestice unutar svakog konstrukta konzistentno mjere istu temeljnu latentnu varijablu. To vrijedi i za konstrukt olakšavajućih uvjeta (FC) ($\alpha = 0,779$; CR (rho_c) = 0,860, CR (rho_a) = 0,815 za cjelokupni uzorak) čija čestica FC4 ima niže vanjsko opterećenje. Prikazani rezultati opravdavaju zadržavanje svih čestica radi bolje usporedivosti s izvornim UTAUT2 modelom i drugim povezanim istraživanjima. Nadalje, konvergentna valjanost mjerena je pokazateljem AVE i vanjskim opterećenjima čestica. AVE vrijednosti za sve konstrukte bile su iznad preporučenog praga od 0,50 (Fornell & Larcker, 1981), što znači da svaki konstrukt objašnjava više od 50 % svojih indikatora i ukazuje na primjerenu konvergentnu valjanost konstrukta. S druge strane, kod vanjskih opterećenja gotovo sve čestice su imale opterećenja iznad preporučenog praga od 0,70 (Hair i sur., 2022), a izuzetak je bila ranije spomenuta čestica FC4 (FC4: *Mogu dobiti pomoć od drugih kada imam poteškoća s korištenjem alata generativne umjetne inteligencije.*). Čestica FC4 na cjelokupnom je uzorku imala vanjsko opterećenje od 0,552 što je ispod idealnog praga. Ipak, nakon provedene provjere unutarnje konzistentnosti konstrukta FC, odluka je autora da se čestica zadrži zbog jasnije usporedivosti dobivenih empirijskih rezultata i nalaza drugih autora. Sličan zaključak navode i Cabero-Almenara i sur., (2025) gdje je čestica FC4 također ispod 0,70. Ipak, Zaid Kilani i sur., (2023) navode kako su oni navedenu česticu izbacili iz mjernog modela jer su vanjska opterećenja spomenute čestice bila ispod preporučenog praga od 0,70. U konačnici, provjerena je i diskriminantna valjanost pomoću Fornell-Larckerovog kriterija i *Heterotrait-Monotrait* (HTMT) omjera. Prema Fornell-Larckerovom kriteriju gdje su na dijagonali prikazane vrijednosti korijena pokazatelja AVE, potvrđuju kako su korelacije među česticama unutar konstrukta snažnije nego korelacije s drugim konstruktima. Ovom provjerom potvrđena je diskriminantna valjanost svih konstrukta u konceptualnom modelu, a jedina uočena malo viša korelacija (0,703 na cjelokupnom uzorku) je bila između očekivanih napora (EE) i olakšavajućih uvjeta (FC), što je konceptualno očekivano jer su konstrukti i teorijski povezani te se oba odnose na lakoću korištenja i dostupnost podrške (Venkatesh i sur., 2003; Shaw i Sergueeva, 2019). Ipak, bitno je naglasiti kako ova vrijednost nije narušila diskriminantnu valjanost. Nadalje, HTMT omjer dodatno je potvrdio diskriminantnu valjanost mjernog modela. Sve HTMT vrijednosti bile su ispod preporučenog praga od 0,90, pa čak i konzervativnijeg praga od 0,85 (Henseler i sur., 2015). Najviša zabilježena HTMT vrijednost bila je 0,823 između očekivanih napora (EE) i olakšavajućih uvjeta (FC), no i ona je i dalje bila u granicama prihvatljivog. Dobiveni rezultati

potvrđuju kako je mjerni model jasno definiran i kako svaki konstrukt mjeri različite aspekte prihvatanja generativne umjetne inteligencije.

Nadalje, u radu je provjerena i kolinearnost među latentnim varijablama koja može predstavljati problem u procjeni jedinstvenog doprinosa svake čestice. U ovom doktorskom radu kolinearnost je provjerena pomoću faktora inflacije varijance (VIF). U okviru provjere na cjelovitom uzorku sve su VIF vrijednosti bile ispod preporučenog praga od 5,0 (Hair i sur., 2021; Kock i Lynn, 2012). Najviša zabilježena VIF vrijednost bila je 4,895 za česticu HM1. Premda visoka, ona ne prelazi navedeni preporučeni prag i ovi rezultati potvrđuju kako multikolinearnost nije problem u ovom mjernom modelu. Ipak, analiza poduzorka s područja Ujedinjenog Kraljevstva ukazuje na povišene VIF vrijednosti čestica HM1 (6,928) i HM2 (5,219) unutar konstrukta hedonističke motivacije. Dobiveni rezultati sugeriraju multikolinearnost između tog konstrukta za ispitanika iz UK-a, što može biti interpretirano kako ispitanici s engleskog govornog područja navedene tvrdnje shvaćaju previše sličnima (HM1 – Korištenje generativne umjetne inteligencije je zabavno; engl. *Using Gen AI is fun*) i HM2 – Korištenje generativne umjetne inteligencije je ugodno; engl. *Using Gen AI is enjoyable*). S obzirom na to da u okviru cjelovitog uzorka VIF vrijednosti navedenih čestica nisu bile problematične, nije bilo potrebe za njihovim izbacivanjem iz mjernog modela.

Premda da je u preliminarnom istraživanju (Biloš i Budimir, 2024) potvrđena visoka korelacija (Pearsonov $r = 0,84$) između namjere ponašanja (BI) i stvarnog korištenja (USE), odnosno previsoka razina multikolinearnosti, jedan od najvećih izazova ovog dokorskog rada bio je unaprjeđenje konstrukta USE. Izvorni UTAUT2 rad koji su predstavili Venkatesh i sur. (2012) ne navodi čestice za konstrukt USE koje bi se univerzalno mogle koristiti u različitim kontekstima istraživanja, stoga su u spomenutom preliminarnom istraživanju korištene čestice po uzoru na Nikolopoulou i sur. (2020). Navedeni izbor nije se pokazao uspješnim te su na temelju analize u poglavlju 4.1.1. odabrane nove tri čestice koje mjere trajanje, učestalost i intenzitet korištenja generativne umjetne inteligencije po uzoru na Venkatesha i sur. (2008). Premda je korelacija između BI i USE i dalje prisutna, što je teorijski i očekivano, PLS-SEM analiza uspjela je razlučiti njihove utjecaje, a svi pokazatelji valjanosti za konstrukt USE bili su zadovoljavajući (kod cjelovitog uzorka: $\alpha = 0,846$; CR (rho_c) = 0,906; CR (rho_a) = 0,862; AVE = 0,764). Dobiveni rezultati sugeriraju kako je odabrani set čestica za konstrukt USE u ovom slučaju metodološki ispravan i prikladan za mjerenje stvarnog korištenja generativne umjetne inteligencije.

8.2. Sinteza glavnih nalaza i usporedba s prethodnim istraživanjima

U ovom dijelu rasprave analiziraju se i interpretiraju rezultati empirijskog istraživanja vezanih za odrednice prihvaćanja i korištenja generativne umjetne inteligencije mjerene proširenim UTAUT2 modelom te se uspoređuju s nalazima relevantne literature na ovu i povezane teme.

Analiza konceptualnog modela na cjelovitom uzorku pokazala je umjereno visoko objašnjenje varijance od 70,1 % za namjeru ponašanja, ($R^2 = 0,701$; $Q^2 = 0,659$) i nešto manje, ali i dalje umjereno snažno objašnjenje varijance od 57,7 % za stvarno korištenje ($R^2 = 0,578$; $Q^2 = 0,546$). Rezultati su približno slični i na poduzorcima, pa tako objašnjenje varijance za BI na tržištu Ujedinjenog Kraljevstva iznosi 73,9 %, a za USE 62,9 %. S druge strane, na hrvatskom tržištu objašnjenje varijance za BI iznosi 72,1 %, dok za USE iznosi 54, 5%. Eksperti za PLS-SEM analize razilaze se u intepretiranju ovih vrijednosti, Hair i sur. (2017a) tumače kako je R^2 preko 0,50 umjereno objašnjena varijanca, a R^2 preko 0,75 visoko objašnjena varijanca, dok Chin (1998) tumači kako je R^2 veći ili jednak od 0,67 visoko objašnjena varijanca, a R^2 veći od 0,33 umjereno objašnjena varijanca. Prethodno navedeni nalazi su u skladu s drugim istraživanjima koji potvrđuju slično objašnjenje varijance za namjeru ponašanja, pa tako Biloš i Budimir (2024) navode da je objašnjenje za BI 65 %, Meet i sur. (2022) navode 69,9 % objašnjenja namjere ponašanja, Salifu i sur., (2024) da je objašnjenje 74 %, Sugumar & Chandra (2021) navode 65 %, Nikolopoulou i sur. (2020) navode 62,7 %, Strzelecki (2023) navodi 73,4 %, Budhathoki i sur. (2024) navode 69,6 %, Baptista i Oliviera (2015) navode 69,1 %, Salifu i sur. (2024) navode 74 %, Gansser i Reich (2021) za svoja tri konteksta navode objašnjenje od 58,3 % do 71,8 % za BI, dok izvorna UTAUT2 teorija navodi vrijednost od 74 % objašnjenja varijance (Venkatesh i sur., 2012). S druge strane, vrijednosti objašnjenja varijance za stvarno korištenje (USE) više variraju s obzirom na to da istraživanja koriste različite čestice za mjerenje korištenja tehnologije, no i dalje su vrijednosti ovog empirijskog istraživanja u skladu s drugim istraživanjima pa tako Baptista i Oliveira (2015) navode objašnjenje varijance za USE od 58,7 %, Oliveira (2016) navodi 61,3 %, Salifu i sur. (2024) navode 52 %, Nikolopoulou i sur. (2020) navode 63,4 %, Chopdar navodi 58,5 % na američkom i 61,8 % na indijskom tržištu, Kasparova (2022) navodi 57,5 %, a izvorni UTAUT2 model navodi objašnjenje varijance kod stvarnog korištenja od 52 % (Venkatesh i sur., 2012).

Promatrajući odrednice koje statistički značajno utječu na namjeru ponašanja (BI), vidljivo je kako su to navika (HT), hedonistička motivacija (HM), povjerenje (TR), očekivani učinak (PE), društveni utjecaj (SI) i osobna inovativnost (PI). S druge strane, kod stvarnog ponašanja (USE)

to su navika (HT), namjera ponašanja (BI) i olakšavajući uvjeti (FC). Jedine odrednice koje nisu statistički značajni prediktori niti jedne zavisne varijable su očekivani napori (EE) i cijena (PV).

Rezultati ovog istraživanja pokazuju kako je očekivani učinak (PE) statistički značajan i umjereno snažan prediktor namjere ponašanja (BI) ($\beta = 0,113$; $p < 0,05$), što je u skladu s osnovnim pretpostavkama UTAUT modela (Venkatesh i sur., 2003), kao i njegovim proširenjem UTAUT2 (Venkatesh i sur., 2012). Ipak, u navedenim se istraživanjima PE pokazuje kao najsnažniji prediktor namjere ponašanja. Slika 13. u potpoglavlju 4.3.1. pokazuje kako je u radovima obuhvaćenima pregledom literature to i najčešće potvrđeni prediktor namjere ponašanja. Više studija iz pregleda literature naglašava kako korisnici žele koristiti napredne tehnologije ako im povećavaju produktivnost i pomažu u zadacima (Gansser & Reich, 2021; Budhathoki i sur., 2024; Grassini i sur., 2024). Zanimljivo je promatrati kako se PE > BI nije potvrdio značajnim u analizi hrvatskog uzorka ($\beta = -0,030$; $p > 0,05$), ali jest u uzorku ispitanika iz Ujedinjenog Kraljevstva ($\beta = 0,316$; $p < 0,05$) i to kao najsnažniji prediktor. Rezultati ukazuju na moguće razlike u motivaciji korisnika iz različitih kulturnih i tržišnih konteksta. U Ujedinjenom Kraljevstvu korisnici očito veću važnost pridaju korisnosti i učinkovitosti tehnologije, dok se u Hrvatskoj ti faktori možda percipiraju slabije ili ih nadjačavaju drugi prediktori namjere ponašanja. Odgovor potencijalno leži u odnosu moderatorskih varijabli koje utječu na ovu vezu. Dodatna analiza na hrvatskom uzorku pokazuje kako PE s uključenim moderatorskim varijablama nije značajan ($\beta = -0,030$; $p > 0,05$), ali postaje značajan prediktor ($\beta = 0,099$; $p < 0,05$) kada se iz modela isključe moderatorske varijable. Nadalje, zanimljivo je za istaknuti kako je PE na uzorku iz UK-a potvrđen kao najsnažniji prediktor namjere ponašanja i u modelu s uključenim moderatorima ($\beta = 0,316$; $p < 0,05$) i u modelu bez moderatora ($\beta = 0,291$; $p < 0,05$).

Očekivani napori (EE), za razliku od očekivanih učinaka (PE), nisu potvrđeni kao statistički značajan prediktor namjere ponašanja ($\beta = -0,021$; $p > 0,05$). Utjecaj EE > BI nije potvrđen ni kod poduzorka iz UK-a ($\beta = -0,041$; $p > 0,05$) niti kod poduzorka iz RH ($\beta = 0,028$; $p > 0,05$). Iz navedenih je rezultata vidljivo kako je promatrani utjecaj EE > BI gotovo beznačajan, a slična je situacija i kod konceptualnog modela bez moderatora ($\beta = -0,052$; $p > 0,05$). Ovakvi nalazi u skladu su s preliminarnim istraživanjem (Biloš i Budimir, 2024), ali i mnogim drugim istraživanjima koji su potvrdili kako EE nije statistički značajan prediktor namjere ponašanja kod alata generativne umjetne inteligencije (Grassini i sur., 2024; Sinaga i sur., 2024; Mohd Rahim i sur., 2022; Salifu i sur., 2024; Lavidas i sur., 2024; Habibi i sur., 2023; Xian, 2021;

Cabrera-Sanchez i sur., 2021), ali i kod drugih tehnologija poput sustava za upravljanje učenjem (Ain i sur., 2016), *fintech* usluga i mobilnog bankarstva (Hassan i sur., 2023; Oliveira i sur., 2016; Baptista i Oliveira, 2015; Thusi i Maduku., 2020) ili internetske trgovine (Shaw i Sergueeva, 2019; Setiyani i sur., 2023). Slika 14. u potpoglavlju 4.3.1. prikazuje kako je utjecaj $EE > BI$ najčešće nepotvrđeni prediktor u ovakvim i srodnim istraživanjima. Ovakvi rezultati sugeriraju da korisnici generativne umjetne inteligencije, vjerojatno zbog relativno jednostavnih i intuitivnih sučelja alata generativne umjetne inteligencije, ne percipiraju da je potrebno uložiti trud u korištenje takvih alata te zbog toga navedena odrednica ne utječe na njihovu namjeru. Menon & Shilpa (2023) sugeriraju da su moderni alati za generativnu umjetnu inteligenciju dizajnirani s naglaskom na intuitivno korisničko sučelje te da se time može objasniti potencijalna beznačajnost očekivanih napora. Određene studije ukazuju na činjenicu kako EE gubi na značajnosti u ovakvim istraživanjima jer je digitalna pismenost među općom populacijom sve veća, a tehnologije su sve intuitivnije za korištenje (Cimperman i sur., 2013; Thusi & Maduku., 2020). Ipak, u modelu poput izvornog UTAUT-a (Venkatesh i sur., 2003) koji ima manje prediktora namjere ponašanja, EE pokazuje značajniji utjecaj na BI (Budhathoki i sur., 2024; Lai i sur., 2024; Sobaih i sur., 2024; Cabero-Almenara i sur., 2025; Xu i sur., 2025).

Društveni utjecaj (SI) još je jedan prediktor koji statistički značajno utječe na namjeru ponašanja (BI) u ovom istraživanju ($\beta = 0,075$; $p < 0,05$). Ovakav nalaz sugerira kako percepcija mišljenja važnih osoba iz okruženja pojedinca utječe na njegovu namjeru da koristi alate generativne umjetne inteligencije. U kontekstu relativno nove tehnologije, pojedinci se često oslanjaju na iskustva i stavove drugih kako bi kreirali vlastite stavove (Venkatesh, 2022). Pregled literature u potpoglavlju 4.3.1. ukazuje na mješovite rezultate za odnos $SI > BI$, pri čemu dio studija pronalazi statistički značajan utjecaj (Sobaih i sur., 2024; Xian, 2021; Budhathoki i sur., 2023), dok drugi ovaj utjecaj ne potvrđuju kao statistički značajan (Garcia de Blanes Sebastian i sur., 2022; Kasparova, 2022; Meet i sur., 2022; Mohd Rahim i sur., 2022; Sinaga i sur., 2024; Lavidas i sur., 2024). Ipak, zanimljivo je za uočiti kako je ovaj prediktor češće potvrđen kao statistički značajan ako je kontekst istraživanja neka telekomunikacijska tehnologija pri čemu pojedinci žele koristiti iste komunikacijske kanale kao i njihovi prijatelji ili obitelji (Ameri i sur., 2020; Nikolopoulou i sur., 2020). Nadalje, na poduzorku iz Ujedinjenog Kraljevstva SI je također potvrđen kao značajan prediktor namjere ponašanja ($\beta = 0,134$; $p < 0,05$) dok na poduzorku ispitanika iz RH navedeni odnos nije potvrđen kao značajan ($\beta = 0,017$; $p > 0,05$). Ovakvi nalazi ukazuju na snažniji društveni pritisak za korištenjem alata generativne umjetne inteligencije na tržištu Ujedinjenog Kraljevstva. Korisnici generativne umjetne

inteligencije u UK-u kao tehnološki iznimno naprednoj ekonomiji, okruženi su digitalno pismenijim društvom, izloženi su pozitivnim primjerima i preporukama za korištenje takve tehnologije zbog čega im je i društveni utjecaj važniji nego kod korisnika iz Hrvatske. Institucionalna podrška, obrazovni programi, vladine inicijative i medijska promocija također doprinose stvaranju percepcije da je korištenje spomenute tehnologije društveno poželjno.

Utjecaj olakšavajućih uvjeta (FC) na namjeru ponašanja (BI) nije bio statistički značajan na cjelokupnom uzorku ($\beta = 0,046$; $p > 0,05$), ali niti na UK i RH tržištima, što je u skladu s nalazima iz pregleda literature gdje se uspoređivanjem analiza iz slike 13. i slike 14. može uočiti kako u promatranim radovima faktor FC češće nije značajan prediktor namjere ponašanja (BI) nego što jest. Ovaj nalaz u skladu je i s preliminarnim istraživanjem (Biloš i Budimir, 2024), ali i mnogim drugim radovima koji istražuju utjecaj olakšavajućih uvjeta na namjeru ponašanja u kontekstu generativne umjetne inteligencije (Maican i sur., 2023; Mohd Rahim i sur., 2022; Strzelecki, 2023; Foroughi i sur., 2024; Grassini i sur., 2024; Lavidas i sur., 2024; Sinaga i sur., 2024; Sobaih i sur., 2024). Ovi rezultati sugeriraju kako dostupnost infrastrukturnih resursa i podrške ne utječe na namjeru da pojedinac koristi alate generativne umjetne inteligencije. Pretpostavka je autora kako razlog leži u činjenici da su alati generativne umjetne inteligencije poput ChatGPT-a ili Geminija lako dostupni, mogu se pokrenuti na gotovo svakom suvremenom uređaju, od pametnih telefona, tableta pa sve do osobnih računala. Osim toga, kako je već spomenuto kod prediktora očekivanih napora, spomenuti alati imaju vrlo intuitivna sučelja i lako ih je naučiti (Strzelecki, 2023) pa zbog toga takva tehnologija rijetko zahtijeva pomoć digitalno spretnijih osoba iz okruženja ili pak korisničke službe (Shaw & Sergueeva, 2019).

Međutim, utjecaj olakšavajućih uvjeta (FC) pokazao se statistički značajnim na stvarno korištenje (USE) ($\beta = 0,086$; $p < 0,05$), što je u skladu s analiziranom literaturom (Cabrera-Sanchez i sur., 2021; Habibi i sur., 2023; Strzelecki, 2023; Lavidas i sur., 2024; Sinaga i sur., 2024). Nadalje, zanimljivo je opažanje kako je utjecaj $FC > USE$ kod poduzorka s UK tržišta snažniji u modelu bez moderatora ($\beta = 0,156$; $p < 0,05$) nego u modelu s uključenim moderacijskim odnosima ($\beta = 0,102$; $p < 0,05$). Također, zanimljivo je uočiti kako utjecaj $FC > BI$ nije značajan, dok utjecaj $FC > USE$ je značajan, a takav je obrazac vidljiv kod cjelovitog uzorka, ali i na poduzorku s tržišta Ujedinjenog Kraljevstva. Prethodna istraživanja pokazuju kako nije rijetko da FC ima statistički značajan utjecaj na USE, ali ne i na BI (Strzelecki, 2023; Ain, 2025; Cabrera-Sanchez i sur., 2021; Nikolopoulou i sur., 2020). Takav obrazac odnosa u modelu može se objasniti činjenicom da dostupnost resursa, tehnička podrška i infrastruktura

nisu važni kod oblikovanja namjere o tome hoće li pojedinac isprobati određenu tehnologiju, osobito ako je jednostavna i lako dostupna, ali postaje jako bitno za stvarno korištenje jer bez pristupa mobitelu, računalu ili internetu korisnik neće biti u stanju kontinuirano koristiti spomenutu tehnologiju. Za ovakav istraživački zaključak dosta je indikativno i kako su Venkatesh i sur. (2003) u izvornom UTAUT-u predložili da FC ima utjecaj samo na USE. Dodatno objašnjenje ovakvom nalazu može biti broj prediktora koje ima zavisna varijabla BI, što znači da druge odrednice poput očekivanih učinaka (PE), društvenih utjecaja (SI), osobne inovativnosti (PI), navike (HT), povjerenja (TR) i hedonističke motivacije (HM) bolje objašnjavaju namjeru ponašanja. Osim toga, ako se promatra organizacijski kontekst, uloga olakšavajućih uvjeta može varirati. Tako u organizacijama s razvijenom infrastrukturom i tehničkom podrškom korisnici mogu imati veću sklonost korištenju alata generativne umjetne inteligencije nego u okruženjima s ograničenijim resursima.

Hedonistička motivacija (HM) statistički je značajan, umjereno snažan i drugi najsnažniji prediktor namjere ponašanja (BI) na ukupnom uzorku ($\beta = 0,239$; $p < 0,05$). Zanimljivo je za uočiti kako je i kod hrvatskog uzorka HM drugi najsnažniji prediktor namjere ponašanja ($\beta = 0,309$; $p < 0,05$). Također, kod cjelovitog uzorka, ali i kod poduzorka s oba tržišta uočava se kako je prediktivna snaga $HM > BI$ snažnija u modelu s uključenim moderatorima pri čemu muškarci više uživaju u korištenju generativne umjetne inteligencije. Ipak, samo se kod cjelovitog uzorka pokazalo kako je spol značajan moderator navedenog odnosa. Dobiveni rezultati u skladu su s drugim istraživanjima u kontekstu generativne umjetne inteligencije i srodnih tehnologija (Strzelecki, 2023; Habibi i sur., 2023; Sinaga i sur., 2024; Mehedi Hasan Emon i sur., 2023; Gansser & Reich, 2021). Faktor PE može se interpretirati kao čimbenik vanjske motivacije jer se usredotočuje na korisnost određene tehnologije (Venkatesh i sur., 2003), dok se HM može okarakterizirati kao čimbenik unutarnje motivacije (Venkatesh i sur., 2012). Visok značaj $HM > BI$ ukazuje da aspekti zabave, užitka i znatiželje igraju važnu ulogu u motiviranju korisnika da koristi tehnologiju (Strzelecki, 2023; Habibi i sur., 2023). Nadalje, Kang i sur. (2025) navode kako hedonistička motivacija ima značajan utjecaj na emocionalnu vrijednost koju korisnici pridaju tehnologiji što posredno oblikuje pozitivne stavove prema korištenju generativne umjetne inteligencije. Bitan je to zaključak jer emocionalne komponente često nadmašuju funkcionalne koristi, posebno kada tehnologija nudi interaktivna i personalizirana iskustva. U skladu s tim, može se zaključiti kako je za namjeru korištenja generativne umjetne inteligencije važno da korisnici uživaju u korištenju spomenute tehnologije što potvrđuje važnost nematerijalnih, intrizičnih motivatora kod korištenja

tehnologije, pogotovo onih tehnologija poput generativne umjetne inteligencije koje nude zanimljive interakcije. Ipak, navedeni nalaz treba promatrati i kroz prizmu efekta noviteta pri čemu korisnici doživljaju povećano zadovoljstvo zbog novosti same tehnologije. S vremenom se korisnici navikavaju na tehnologiju i efekt može padati (Wells i sur., 2010; Hopp & Gangadharbatla, 2016).

Vrijednost cijene ili jednostavnije cijena (PV) nije statistički značajan prediktor namjere ponašanja na cjelovitom uzorku ($\beta = 0,023$; $p > 0,05$) niti je na poduzorcima iz RH i UK. Navedeni utjecaj nije potvrđen niti u modelima bez moderatora. Ovaj nalaz nije u skladu s izvornim UTAUT2 modelom (Venkatesh i sur., 2012), no s obzirom na dominantnu prisutnost besplatnih verzija alata generativne umjetne inteligencije, u određenoj mjeri je i očekivano da faktor cijene nema snažan utjecaj na namjeru isprobavanja i korištenja navedenih alata (Law, 2024; Strzelecki, 2023, Biloš & Budimir, 2024; Faraon i sur., 2025; Ali i sur., 2024). Među ispitanicima provedenog istraživanja 77 % ispitanika ne koristi plaćenu verziju niti jednog alata generativne umjetne inteligencije, odnosno samo 23 % njih koristi. Druga istraživanja koja se također bave tehnologijama čije su osnovne funkcionalnosti besplatne, poput mobilnog bankarstva, također zaključuju kako cijena nije prediktor namjere ponašanja (Thusi & Maduku., 2020; Baptista & Oliveira, 2015; Oliveira i sur., 2016). Strzelecki (2023) je u svom konceptualnom modelu izostavio faktor cijene uz napomenu kako navedeni alat generativne umjetne inteligencije većina korisnika koristi u besplatnoj formi. S obzirom na to da ne postoji monetarni trošak kao barijera u isprobavanju ove tehnologije, on ne utječe na namjeru korisnika i zbog toga je ova hipoteza odbačena uz naglasak kako kod nje nema ni moderatorskih utjecaja.

Navika (HT) je još jedna od varijabli čiji se utjecaj unutar konceptualnog modela mjeri i na namjeru ponašanja (BI) i stvarno korištenje (USE). Ovo istraživanje potvrđuje kako je navika najснаžniji prediktor namjere ponašanja ($\beta = 0,336$; $p < 0,05$) i stvarnog korištenja ($\beta = 0,409$; $p < 0,05$). U analizi odnosa iz pregledene literature, u okviru potpoglavlja 4.3.1., vidljivo je kako većina istraživanja potvrđuje tvrdnju da je navika značajan prediktor namjere ponašanja i stvarnog korištenja, a nerijetko je i najснаžniji prediktor namjere ponašanja (Baptista i Oliveira, 2015; Biloš i Budimir, 2024; Nikolopoulou i sur., 2020; Merhi i sur., 2019; Hew i sur., 2015; Khan i sur., 2022; Sinaga i sur., 2024), ali i stvarnog korištenja (Tamilmani i sur., 2020; Kasparova, 2022; Lavidas i sur., 2024; Baptista i Oliveira, 2015). U pravilu, najснаžniji prediktori namjere ponašanja su PE i HT, dok su kod stvarnog korištenja to BI i HT (Faraon i sur., 2025). Izrazita snaga navike kao prediktora sugerira da korištenje alata generativne umjetne inteligencije danas postaje dio rutine, postaje nešto sasvim normalno u svakodnevicu

pojedince (Gansser i Reich, 2021). Drugim riječima, pojedinci koji su integrirali korištenje alata generativne umjetne inteligencije u svakodnevne zadatke poput pisanja e-pošte, generiranje ideja, traženje informacija ili prijevoda teksta, to rade više iz navike nego zbog nekakve provjere učinkovitosti alata, njihove jednostavnosti ili užitka korištenja. U poduzorku iz Ujedinjenog Kraljevstva HT nije najsnažniji prediktor namjere ponašanja, već je PE, a nije ni najsnažniji prediktor stvarnog korištenja, već je BI. Ovo potvrđuje i prethodno spomenutu teoriju kako se navedeni prediktori najčešće ističu kao najsnažniji prediktori. Činjenica da je navika jedan od najsnažnijih prediktora namjere ponašanja, ali i stvarnog korištenja, može se protumačiti i petljom navika (engl. *Habit loop*) jer ispitanici izražavaju kako su im i faktori unutarnje (HM, PI), ali i vanjske motivacije (PE, SI, TR) značajni za namjeru korištenja. S obzirom na to da postoji snažan motiv koji korištenje generativne umjetne inteligencije čini korisnim, društveno poželjnim i zanimljivim, pojedinci stvaraju naviku korištenja (Duhigg, 2013).

Osobna inovativnost (PI) kao konstrukt kojim je proširen izvorni UTAUT2 u okviru ovog konceptualnog modela potvrđen je kao statistički značajan prediktor namjere ponašanja ($\beta = 0,070$; $p \leq 0,05$). I premda je p vrijednost u okviru cjelovitog uzorka rubna $p = 0,050$, ona se prihvatila kao statistički značajna jer je t vrijednost blago iznad 1,960, točnije 1,963. Ono što je iz rezultata testiranja odnosa u modelu bez moderatora vidljivo jest porast prediktivne snage $PI > BI$ ($\beta = 0,121$; $t = 4,523$; $p = 0,000$) što znači da su pretpostavljeni moderatori dob i spol zapravo negativno utjecali na prediktivnu snagu osobne inovativnosti kao čimbenika u namjeri prihvaćanja alata generativne umjetne inteligencije. Isti obrazac uočen je i kod poduzoraka s oba tržišta, uz napomenu kako na hrvatskom uzorku $PI > BI$ nije bio statistički značajan u modelu s moderatorima, ali jest u testiranju odnosa modela bez moderatora. Dobiveni rezultati u skladu su s postojećim nalazima da su inovativniji pojedinci skloniji usvajanju novih tehnologija (Agarwal i Prasad, 1998; Gansser i Reich, 2021; Garcia de Blanes Sebastian i sur., 2022; Strzelecki i sur., 2023), a što generativna umjetna inteligencija definitivno jest. Ipak, zanimljiv je nalaz koji navode Faraon i sur. (2025) u kojem uspoređuju nordijske zemlje i SAD, a u kojem potvrđuju kako $PI > BI$ nije statistički značajan u nordijskim zemljama, ali jest u SAD-u kao kulturološkom okruženju koje naglašava individualizam. Značaj osobne inovativnosti sugerira sklonost pojedinaca ka tehnološkim novitetima i eksperimentiranju, a što ima važnu ulogu u namjeri korisnika da isproba novu tehnologiju. Očito je kako u tehnološki naprednijim ekonomijama ispitanici imaju snažniju želju za isprobavanjem novih tehnologija.

Još jedan konstrukt s kojim je proširen konceptualni model ovog doktorskog rada je povjerenje (TR) čiji se utjecaj mjeri na zavisnu varijablu namjere ponašanja (BI). Provedeno empirijsko istraživanje potvrdilo je kako povjerenje umjereno snažno objašnjava namjeru pojedinca da koristi alate generativne umjetne inteligencije ($\beta = 0,187$; $p < 0,05$). Ovo je iznimno važan empirijski nalaz s obzirom na prirodu alata generativne umjetne inteligencije gdje pitanja pouzdanosti, točnosti, etičnosti i privatnosti često dolaze u prvi plan (Kirova i sur., 2023). Također, dobiveni rezultati u skladu su s prethodnim istraživanjima koja su uključivala čimbenik povjerenja (Cabrera-Sanchez i sur., 2021; Garcia de Blanes Sebastian i sur., 2022; Kilani i sur., 2023; Mohd Rahim i sur., 2022; Salifu i sur., 2024; Vimalkumar i sur., 2021; Lai i sur., 2024; Mehedi Hasan Emon i sur., 2023; Merhi i sur., 2019; Ali i sur., 2024). Relativno visok značaj povjerenja ukazuje da pojedinci moraju vjerovati u sposobnost generativne umjetne inteligencije da pruži točan i pouzdan sadržaj, uz uvjet da će njihovi privatni podatci ostati sigurni kako bi im pružili šansu i probali koristiti određeni alat. Moderator i navedenog odnosa gotovo su potpuno beznačajni te je utjecaj $TR > BI$ približno jednak i u modelu s moderatorima, ali i u modelu bez moderatora.

Kao što teorija nalaže (Sheppard i sur., 1988; Venkatesh i sur., 2003; Venkatesh i sur., 2012), namjera ponašanja (BI) predviđa stvarno korištenje tehnologije (USE) i to je ovim empirijskim istraživanjem dokazano ($\beta = 0,310$; $p < 0,05$). Premda čestice konstrukta USE nisu usklađene s izvornim UTAUT2 modelom (Venkatesh i sur., 2012), već modificirane po uzoru na Venkatesh i sur. (2008), rezultati pokazuju kako su one metodološki valjane i ispravno integrirane u model. Značajan odnos $BI > USE$ potvrđen je u cjelovitom uzorku, ali i kod oba poduzorka (RH i UK). Odnos je dodatno testiran i na modelu bez moderatora koji je potvrdio kako moderirajući utjecaj iskustva zapravo smanjuje prediktivnu moć BI na USE jer je kod cjelovitog uzorka, ali i oba poduzorka, potvrđeno kako je utjecaj $BI > USE$ snažniji u modelu bez moderatora.

Za bolje razumijevanja moderatora, važno je spomenuti kako više stručnih istraživanja iz ovog područja navodi kako generativnu umjetnu inteligenciju češće koriste muškarci nego žene, ali i mlađi češće nego stariji (Bick i sur., 2024; Deloitte, 2024; Salesforce, 2024; Lin & Parker, 2025; Aldoroso i sur., 2024a, Microsoft, 2024). Ipak, navedena istraživanja mjere postotak korisnika među općom populacijom, dok se empirijsko istraživanje u okviru ovog doktorskog rada usredotočuje samo korisnike generativne umjetne inteligencije, a ne opću populaciju.

Analiza modela bez uključenih moderatora pokazala je približne rezultate kao i model s moderatorima, no ipak određenih promjena ima i vrlo su zanimljive za kontekst ovog doktorskog rada. Objašnjenje varijance (R^2) bila je niža, kako za BI (67,9 % nasuprot 70,1 % u

modelu s uključenim moderatorima), tako i za USE (54,2 % u odnosu na 57,7 % u modelu s uključenim moderatorima). Premda se moderatori u ovom istraživanju u većoj mjeri nisu pokazali statistički značajnima, trebati uzeti u obzir kako je moderirajućih efekata u konceptualnom modelu prilično puno (čak 30) i premda pojedinačno nisu značajni, kumulativno mogu doprinijeti ukupnom objašnjenju modela. Osim toga, zanimljivo je za uočiti kako utjecaj $PE > BI$ i $PI > BI$ raste u odsustvu moderatora. Ovakav rezultat sugerira kako su moderatori umanjivali prediktivnu snagu navedenih latentnih varijabli. S druge strane, EE, FC i PV ostali su neznačajni prediktori namjere ponašanja i u modelu bez moderatora. Ovo dodatno potvrđuje njihovu manju relevantnost u formiranju namjere pojedinca da koristi alate generativne umjetne inteligencije. Ipak, prediktoru FC u ovom dijelu treba posvetiti dodatnu pozornost s obzirom na to da je u modelu bez moderatora kod uzorka iz UK-a ipak pokazao statističku značajnost, ali negativnu. Ovo istraživanje nije izuzetak, mnogobrojna istraživanja pokazuju kako se tek poneki moderirajući efekt pokaže značajnim, a u nekima niti jedan. Tako primjerice, istraživanje o prihvaćanju i korištenju umjetne inteligencije među kineskim studentima navodi kako dob, spol i iskustvo ne moderiraju niti jedan odnos u UTAUT2 modelu (Xu i sur., 2024). Grassini i sur. (2024) navode sličan istraživački nalaz koji kaže kako dob, spol i iskustvo ne moderiraju niti jedan odnos u istraživanju prihvaćanja i korištenja generativne umjetne inteligencije među norveškim studentima. Lavidas i sur. (2024) navode iste empirijske nalaze na istraživanju među grčkim studentima, Biloš i Budimir (2024) među hrvatskim studentima, dok Strzelecki (2023) isti zaključak navodi kod istraživanja među poljskim studentima. Huang i sur. (2024) navode kako dob i spol nisu značajni moderatori u modelu koji mjeri stavove prema prihvaćanju umjetne inteligencije u zdravstvu. Ipak, Korkmaz i sur. (2022) navode kako dob i spol značajno moderiraju samo utjecaj hedonističke motivacije na namjeru ponašanja pri čemu je utjecaj jači kod žena, ali ističe i utjecaj povjerenja na ponašanje pri čemu je utjecaj jači kod muškaraca. Ovakvi empirijski nalazi upućuju na to da demografske razlike među ispitanicima i njihovo prethodno iskustvo korištenja tehnologije sve manje utječu na prihvaćanje i korištenje suvremenih tehnologija, posebno jednostavnih, a korisnih i zabavnih tehnologija poput generativne umjetne inteligencije.

8.3. Implikacije istraživanja

8.3.1. Teorijske implikacije istraživanja

Teorijske implikacije ovog istraživanja doprinose razumijevanju prihvaćanja i korištenja alata generativne umjetne inteligencije na nekoliko načina. Prvenstveno, potvrđena je teorijska

valjanost UTAUT2 modela (Venkatesh i sur., 2012), razvijenog kao okvir za istraživanja različitih konteksta iz područja prihvaćanja tehnologije. UTAUT2 jedan je od najutjecajnijih okvira za razumijevanje prihvaćanja tehnologije. U ovom radu analiziran je razvoj modela UTAUT te njegova evolucija u model UTAUT2. Shodno tomu, jedan od ključnih doprinosa je testiranje postojećih konstrukta kao što su očekivani učinak (PE), očekivani naponi (EE), društveni utjecaj (SI), olakšavajući uvjeti (FC), hedonistička motivacija (HM), vrijednost cijene (PV) i navika (HT), ali i identificirana potreba za uvođenjem dodatnih konstrukta. Na temelju pregleda literature prepoznata je važnost osobne inovativnosti (PI) i povjerenja (TR) kao značajnih čimbenika u prihvaćanju generativne umjetne inteligencije. Teorijski okvir koji su predstavili Venkatesh i sur. (2012) vrlo je fleksibilan i omogućuje istraživačima da konceptualni model prilagode kontekstu istraživanja, a vođenjem navedenih konstrukta PI i TR, ovaj doktorski rad teorijski obogaćuje UTAUT2 kao jedan od najvažnijih teorijskih okvira za prihvaćanje tehnologije i prilagođava ga istraživanjima u kontekstu generativne umjetne inteligencije. Nadalje, istraživanjem je potvrđeno kako očekivani naponi (EE) i cijena (PV) ne funkcioniraju kao prediktori namjere ili stvarnog korištenja što je također važan empirijski doprinos ovog dokorskog rada.

Konceptualni model ovog istraživanja koristio je i modificirane čestice konstrukta korištenja (USE). S obzirom na to da se Venkatesh i sur. (2012) pri izradi UTAUT2 modela usredotočuju na mjerenje korištenja mobilnog interneta, njihove čestice za mjerenje korištenja nisu univerzalno primjenjive na širok spektar tehnologija, pa shodno tome niti na istraživanja u kontekstu generativne umjetne inteligencije. Ovaj rad imao je za cilj prilagođavanje čestica konstrukta USE koje će biti primjenjive i relevantne za veći spektar istraživanja u kontekstu generativne umjetne inteligencije i srodnih tehnologija. Dobiveni rezultati ukazuju na to da su čestice konstrukta USE u ovom radu, skrojene po uzoru na Venkatesha i sur. (2008), valjane i dobro prilagođene kontekstu istraživanja te da se mogu koristiti za mnoga daljnja istraživanja na ovu i srodne teme, a što predstavlja važan teorijski doprinos ovog dokorskog rada.

Dodatni metodološki i empirijski doprinos leži u usporednoj analizi prihvaćanja i korištenja generativne umjetne inteligencije u dva različita ekonomska i kulturološka aspekta, u Republici Hrvatskoj kao ekonomiji u razvoju i Ujedinjenom Kraljevstvu kao naprednoj ekonomiji po IMF-ovom Indeksu spremnosti na umjetnu inteligenciju (IMF, 2024).

Nalazi o značajnosti, odnosno beznačajnosti moderatora u istraživanju u većoj mjeri ukazuju na činjenicu da su alati generativne umjetne inteligencije gotovo sveprisutni te da je utjecaj tradicionalnih demografskih moderatora možda slabiji nego kod nekih drugih tehnologija. Alati

generativne umjetne inteligencije u većoj su mjeri vrlo jednostavni i intuitivni i gotovo svi s osnovnom razinom digitalne pismenosti mogu koristiti ovu tehnologiju. Iz navedenog proizlazi da je tradicionalne demografske moderatore možda potrebno zamijeniti nekim specifičnim moderatorima koji su više prilagođeni kontekstu istraživanja, kao primjerice razina digitalne pismenosti.

Ukratko, teorijski doprinosi ovog doktorskog rada su višeslojni. Počevši s činjenicom da je UTAUT2 model uspješno proširen s konstruktima PI i TR koji ga prilagođavaju istraživanjima u kontekstu generativne umjetne inteligencije te povećavaju objašnjenje modela. Nadalje, snažan i važan doprinos ovog istraživanja je usavršavanje mjernih čestica konstrukta USE. Naravno, osim navedenih promjena u modelu, važno je spomenuti i dobivene rezultate iz postojećeg okvira UTAUT2 modela koji omogućuju bolje razumijevanje odrednica koje utječu na prihvatanje i korištenje generativne umjetne inteligencije. U konačnici, svi navedeni faktori testirani su na dva tržišta te se navedene razlike među tržištima također mogu tumačiti kao važan teorijski doprinos ovog doktorskog rada.

8.3.2. Praktične implikacije istraživanja

Praktični ili aplikativni doprinos ovog doktorskog rada prvenstveno se očituje u pružanju konkretnih, praktičnih uvida i preporuka različitim dionicima zainteresiranima za prihvatanje i korištenje generativne umjetne inteligencije, a među kojima treba izdvojiti marketinške stručnjake, stručnjake za ljudske resurse, obrazovne institucije, tvorce zakonodavnog okvira iz domene umjetne inteligencije, ali i poslovne subjekte te pojedince čiji je cilj bolje razumijevanje korisnika generativne umjetne inteligencije.

Nadalje, važan aplikativni doprinos je objašnjenje i formulacija rezultata istraživanja s naglaskom na prediktore namjere korištenja i stvarno korištenje generativne umjetne inteligencije. Identificiranjem tih ključnih odrednica praktičarima je prezentiran smjer u kojem trebaju ulagati svoje napore. Primjerice, očito je kako korisnici podrazumijevaju da su alati generativne umjetne inteligencije jednostavni za korištenje te da ne moraju ulagati dodatni napor i vrijeme za privikavanje na alat. Ipak, rezultati istraživanja ukazuju na to da postojeća tehnička infrastruktura i dostupna podrška također nisu bitni za kreiranje namjere korištenja, ali jesu za nastavak korištenja.

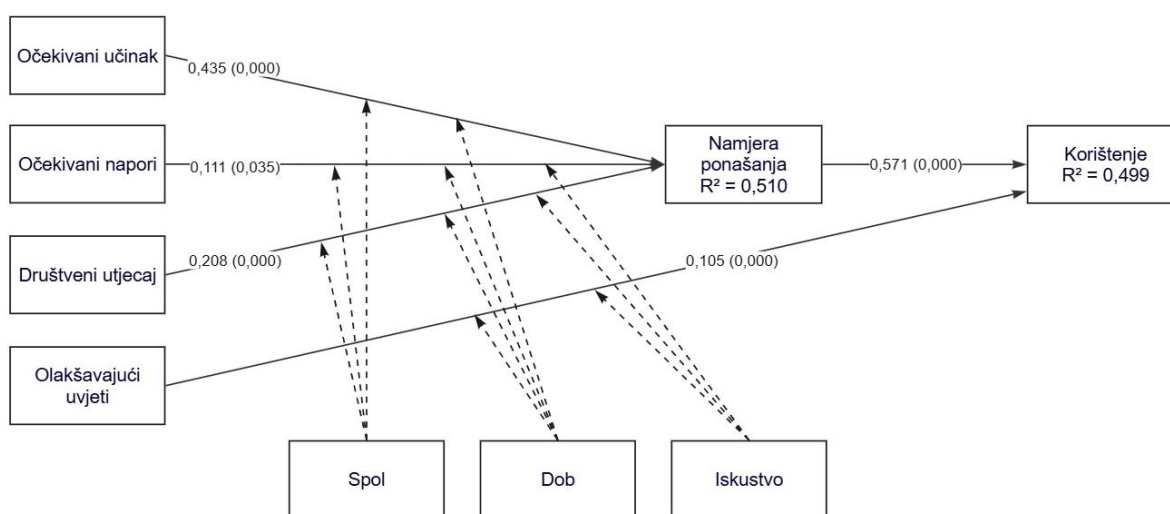
Posebno vrijedan aplikativni doprinos je analiza i ukazivanje na različitosti u namjeri korištenja i stvarnom korištenju generativne umjetne inteligencije među različitim tržišnim segmentima, s posebnim naglaskom na razlike između ispitanika iz Ujedinjenog Kraljevstva i hrvatskih

ispitanika. Usporedna analiza između Republike Hrvatske i Ujedinjenog Kraljevstva, odnosno ekonomije u razvoju i napredne ekonomije po IMF-ovom Indeksu spremnosti na umjetnu inteligenciju (IMF, 2024), nudi važan aplikativni doprinos za globalne tvrtke i međunarodne politike. Primjerice, u Hrvatskoj kao zemlji u razvoju po pitanju umjetne inteligencije, naglasak je na navikama korisnika, hedonističkoj motivaciji i povjerenju kao ključnim odrednicama prihvaćanja alata generativne umjetne inteligencije. Vanjski faktori poput društvenog utjecaja ili razine osobne inovativnosti u usporedbi sa svojim okruženjem imaju značajno manji utjecaj. Političke i obrazovne institucije u Republici Hrvatskoj mogle bi se usmjeriti na povećanje digitalne pismenosti i boljeg poznavanja alata generativne umjetne inteligencije te promicanje istih kao društveno prihvatljivih kako bi očekivani učinak (PE), društveni utjecaj (SI) i osobna inovativnost (PI) postali značajniji pokretači namjere korištenja.

Rezultati empirijskog istraživanja daju jasan uvid u identifikaciju ključnih odrednica za poticanje namjere korištenja alata generativne umjetne inteligencije. S obzirom na to da se hedonistička motivacija, navika i povjerenje ističu kao snažni prediktori, tvrtke koje razvijaju i promoviraju alate generativne umjetne inteligencije trebaju se usmjeriti na naglašavanje zabave i ugođe kod korištenja takvih alata, na poticanje stvaranja navika i izgradnju povjerenja. Ključne su to komponente koje bi mogle privući pasivne korisnike i nekorisnike generativne umjetne inteligencije. Ipak, one nisu jedine važne komponente, bitno je osigurati da korištenje takve tehnologije postane društveno poželjno, a ne stigmatizirano kao nešto negativno ili nemoralno u poslovnom svijetu. Dapače, hrvatsko tržište u ovom je smislu u lošijem položaju nego tržište UK-a jer je značajno manje edukacijsko-informativnog sadržaja o alatima umjetne inteligencije dostupno na hrvatskom jeziku, pa shodno tomu imaju i manje prilike učiti iz dostupnih edukacijskih sadržaja na internetu, a istovremeno su uskraćeni za informacije o mogućnostima koje integracija generativne umjetne inteligencije nosi sa sobom. Dakle, sinergija državnih i obrazovnih institucija s tvrtkama iz privatnog sektora nužna je za popularizaciju generativne umjetne inteligencije, njezinih mogućnosti, prednosti, ali i rizika. Time bi se stvorila snažnija potreba tržišta za novim edukacijskim programima iz ovog polja, korištenje generativne umjetne inteligencije postalo bi društveno poželjnije, a posljedično bi korisnici dobili naviku korištenja takve tehnologije, koja je istovremeno zabavna za korištenje, ali i snažno poboljšava učinkovitost pojedinca.

8.4. Usporedba konceptualnog modela s prethodnicima

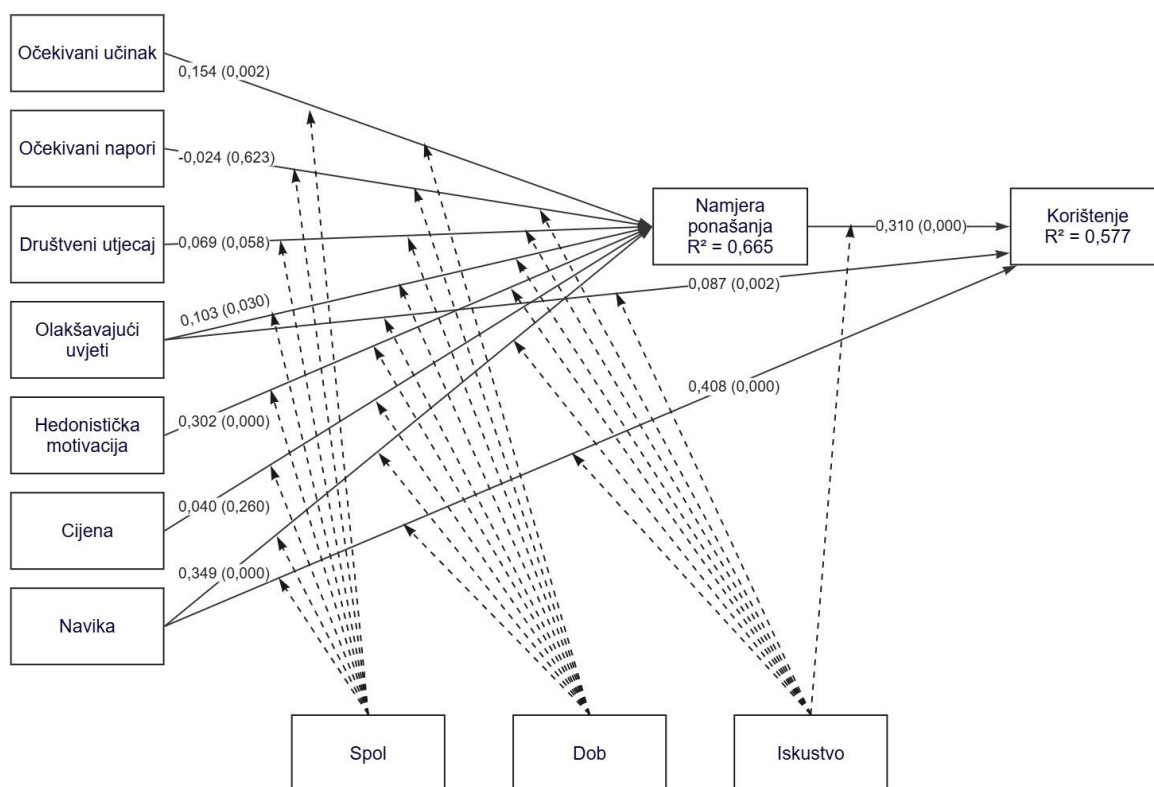
Slika 21., ispod teksta, prikazuje rezultate ovog empirijskog istraživanja na UTAUT modelu, izvornom teorijskom okviru iz 2003. godine. Ipak, treba uzeti u obzir kako u odnosu na izvorni UTAUT model (Venkatesh i sur., 2003) nedostaje moderator dobrovoljnosti. Vidljivo je kako je objašnjenje varijance kod namjere ponašanja tek 51 % ($R^2 = 0,510$), dok je objašnjenje varijance kod stvarnog korištenja 49,9 % ($R^2 = 0,499$) što su relativno niske vrijednosti kod tumačenja objašnjenja modela. Nadalje, vidljivo je kako u modelu sve nezavisne varijable statistički značajno utječu na zavisne varijable, a od moderatora u modelu, potvrdilo se samo kako spol moderira odnos između očekivanog učinka (PE) i namjere ponašanja (BI).



Slika 21. Empirijski rezultati na UTAUT modelu

Izvor: prilagođeno prema Venkateshu i sur. (2003)

S druge strane, na slici 22. vidljivo je kako je UTAUT2 u odnosu na izvorni UTAUT proširen s još tri konstrukta, a to su hedonistička motivacija (HM), cijena (PV) i navika (HT). U provedenoj analizi utvrđeno je kako su očekivani učinak (PE), olakšavajući uvjeti (FC), hedonistička motivacija (HM) i navika (HT) statistički značajni prediktori namjere ponašanja (BI). Ipak, rezultati analize ukazuju i na činjenicu kako očekivani napori (EE), društveni utjecaj (SI) i cijena (PV) nisu značajni prediktori namjere ponašanja. Nadalje, olakšavajući uvjeti (FC), navika (HT) i namjera ponašanja (BI) potvrđeni su kao značajni prediktori stvarnog korištenja (USE). Jedini moderirajući efekt u modelu koji se pokazao značajnim je utjecaj spola na odnos između hedonističke motivacije (HM) i namjere ponašanja (BI). Objašnjenje varijance kod namjere ponašanja je 66,5 % ($R^2 = 0,665$), dok je kod korištenja 57,7 % ($R^2 = 0,577$). Vidljivo je kako je objašnjenje varijance značajno veća u odnosu na izvorni UTAUT.



Slika 22. Empirijski rezultati na UTAUT2 modelu

Izvor: prilagođeno prema Venkateshu i sur. (2012)

Ovakvi nalazi u skladu su s teorijskom osnovom koju navode Venkatesh i sur., (2003) te Venkatesh i sur. (2012). U predstavljanju UTAUT2 modela Venkatesh i sur. (2012) navode kako se objašnjenje varijance kod novog modela podiže s 56 % na 74 % (rast od 18 postotnih poena) dok se za objašnjenje varijance podiže s 40 % na 52 % (rast od 12 postotnih poena). Dakle, za namjeru ponašanja prediktivna snaga UTAUT2 modela jača je 15,5 postotnih poena, dok je za stvarno korištenje (USE) prediktivna snaga modela veća za 7,9 postotnih poena. Uočeni rast kod istraživanja u kontekstu generativne umjetne inteligencije nešto je manji, ali i dalje vrlo značajan.

Zaključno se može reći kako je konceptualni model ovoga rada unaprijedio postojeći teorijski okvir koji su predstavili Venkatesh i sur. (2012), valjano ga prilagodio istraživanju u kontekstu generativne umjetne inteligencije dodajući mu konstrukte osobne inovativnosti i povjerenja što je rezultiralo porastom objašnjenja varijance s 66,5 % na 70,1 % (slika 18.). Ipak, shodno empirijskim nalazima iz ovog istraživanja, ali i nalaza navedenih u poglavlju 4., vidljiv je prostor za dodatna unaprjeđenja konceptualnog modela.

8.5. Ograničenja rada

Premda je istraživanjem obuhvaćen kvotni uzorak s ciljem što reprezentativnijeg uzorka, i dalje se može uočiti kako kod ispitanika iz Republike Hrvatske ima prekomjerno zastupljenog visoko obrazovanog stanovništva, čak 2/3 ispitanika, od čega čak 11,5 % ispitanika ima završen poslijediplomski studij. Iako to nije reprezentativan uzorak opće populacije (DZS, 2023), studije pokazuju kako visokoobrazovani češće koriste alate generativne umjetne inteligencije (Lin & Parker, 2025; Aldaroso i sur., 2024a, Microsoft, 2024; Bick i sur., 2024). Situacija s postotkom visokoobrazovanih je relativno slična i na ispitanicima iz Ujedinjenog Kraljevstva premda su u tom poduzorku realnije zastupljeni ispitanici sa završenim poslijediplomskim studijem. Ipak, kvotnim uzorkom šire populacije u Hrvatskoj i Ujedinjenom Kraljevstvu pokušala su se izbjeći ograničenja zbog manjka generalizabilnosti dobivenih rezultata na koja upozoravaju drugi autori. Riječ je o ograničenjima poput usmjeravanja na preusko geografsko područje (Strzelecki, 2023; Habibi i sur., 2023) ili na previše specifičnu demografsku skupinu (Grassini i sur., 2024; Foroughi i sur., 2024, Strzelecki, 2023).

S metodološke strane, modeli o prihvaćanju tehnologija poput UTAUT2 modela, pa čak i u ovakvoj proširenoj formi, možda ne obuhvaćaju sve relevantne faktore koji utječu na prihvaćanje i korištenje tehnologija poput generativne umjetne inteligencije. Različita istraživanja potvrđuju kako faktori poput anksioznosti (Budhathoki i sur., 2024), straha od tehnologije, (Leschanowsky i sur., 2024), održivosti (Gansser i Reich, 2021), percipirane privatnosti (Merhi i sur., 2019), kvalitete sadržaja (Tantra i Ariyanti, 2017) i znatiželje (Sinaga i sur., 2024) mogu imati značajnu ulogu u predikciji namjere ponašanja i stvarnog korištenja generativne umjetne inteligencije.

Nadalje, generička priroda određenih konstrukta u modelu može dovesti do različitih interpretacija kod ispitanika, a zbog čega može doći i do konceptualnog poklapanja. Ispitanici pilot istraživanja konstatirali su kako postoji puno sličnih tvrdnji koje ih zbunjuju te da imaju osjećaju kao da su već odgovorili na to isto pitanje. Ispitanici pilot istraživanja nadodaju kako im pojedini elementi pitanja nisu posve jasni, kao primjerice tko su „osobe koje su mi važne, a koje smatraju kako trebam koristiti alate generativne umjetne inteligencije“, jesu li to prijatelji, roditelji, profesori, šef ili netko treći? Takve nejasnoće, ali i sličnosti pojedinih konstrukta kao što su očekivani naponi (EE) i olakšavajući uvjeti (FC) mogu narušiti preciznost modeliranja. Ipak, unutar ovog istraživanja to nije bio slučaj, pokazatelji unutarnje konzistentnosti i konvergentne valjanosti potvrđuju kako su podatci mjernog modela valjani. Izuzetak je

faktorsko opterećenje čestice FC4 koje je ostavljeno u modelu jer ne narušava stabilnost modela, a doprinosi boljoj usporedivosti s drugim istraživanjima. Na podacima poduzorka s područja Ujedinjenog Kraljevstva zapažena su malo povišena odstupanja VIF vrijednosti kod čestica HM1 i HM2 što ukazuje na dozu multikolinearnost među česticama konstrukta hedonističke motivacije (HM). Ipak, s obzirom na to da navedeni pokazatelji nisu povišeni na razini cjelokupnog uzorka, navedena se odstupanja ne tumače kao problematična i čestice nisu izbačene iz mjernog modela.

Još jedno ograničenje istraživanja koja se usredotočuju na prihvatanje tehnologije mjerene UTAUT2 modelom je u načinu mjerenja konstrukta USE (korištenja). Riječ je o bitnom metodološkom ograničenju u proučavanju korištenja tehnologije jer različite studije provode različite pristupe u mjerenju tog konstrukta. Neke studije USE tretiraju kao reflektivni konstrukt, mjereno s više čestica koja bi trebale korelirati i održavati latentnu varijablu (Nikolopoulou i sur., 2020; Lavidas i sur., 2024) i kod takvih istraživanja Cronbachova alfa, kompozitna pouzdanost i AVE pokazatelj navedenog konstrukta mora biti zadovoljavajuće razine. U ovom doktorskom radu korišten je pristup mjerenja korištenja kao reflektivnog konstrukta. S druge strane, neka istraživanja koriste USE kao formativni konstrukt gdje se ne očekuje da indikatori visoko koreliraju već da se smatraju uzrocima konstrukta USE (Gansser & Reich, 2021). Drugi problem s mjerenjem korištenja jest relativno niska vrijednost objašnjenja varijance (R^2). Prediktori korišteni uključeni u UTAUT2 model, namjera ponašanja (BI), olakšavajući uvjeti (FC) i navika (HT), tek djelomično objašnjavaju dio ukupne varijance u stvarnom ponašanju korisnika što sugerira da možda postoje drugi značajni, neistraženi i neizmjereni faktori koji utječu na stvarno korištenje generativne umjetne inteligencije. Longitudinalno ili dvoetažno istraživanje u kojem bi se stvarno korištenje pratilo kroz vrijeme pružilo bi jasniji uvid u stvarne pokazatelje korištenja, ali bi i rasvijetlilo promjene u percepciji generativne umjetne inteligencije. Ipak, longitudinalna istraživanja sama po sebi imaju svoje nedostatke, a to su prije svega složenost takvog istraživanja, financijski troškovi i osipanje uzorka. Nadalje, još jedno ograničenje kod mjerenja konstrukta USE, ali i drugih konstrukta je priroda upitnika kojim se mjeri samoprocjena ispitanika (engl. *self-reporting*). Riječ je o standardnom metodološkom ograničenju istraživanja ovakvog karaktera, a nedostatak mu je potencijalna pristranost ispitanika, netočnost prisjećanja, pristranost društveno poželjnog, a ne stvarnog ponašanja, ali i nedostatak precizne svijest o vlastitim navikama i ponašanju.

Jedno vrlo specifično ograničenje ovoga rada je generaliziranje pojma generativne umjetne inteligencije, odnosno alata za korištenje iste. Naime, u uvodnom dijelu istraživačkog

instrumenta ispitanike se pita koje sve alate generativne umjetne inteligencije koriste i tamo je nabrojan 21 alat uz mogućnost da se dodatno nadopíše alat koji možda nije naveden. Iako je uvjerljivo najkorišteniji i najpopularniji alat generativne umjetne inteligencije ChatGPT, to nije istoznačnica s generativnom umjetnom inteligencijom. U nastavku istraživačkog instrumenta ispitanike se pita o tome koliko im je teško naučiti koristiti alat, koliko im alat pomaže u produktivnosti, koliko im je alat zanimljiv, kakav mu je omjer cijene i kvalitete ili koliko mu vjeruju, a ispitanici pri tome mogu razmišljati o različitim alatima ili o nekom specifičnom alatu, a u konačnici su njihovi rezultati generalizirani kao da su svi razmišljali o istome. Ipak, ovakvo istraživanje prisutno je i kod istraživanja drugih tehnologija gdje naglasak nije na jednom sustavu, proizvođaču proizvoda ili pružatelju usluge, kao primjerice kod sustava za e-učenje (Air i sur., 2015; Meet i sur., 2022), uređaja za pametnu kuću (Baudier i sur., 2018; Gothesen i sur., 2023), autonomnih vozila (Nordhoff i sur., 2022; Bellet & Banet, 2023), glasovnih asistenata (Choudhary i sur., 2024; Vimalkumar i sur., 2021), mobilnog bankarstva i drugih *fintech* usluga (Thusi i Maduku, 2020; Hassan i sur., 2023; Kilani i sur., 2023; Merhi i sur., 2019; Baptista i Oliveira, 2015). Autor je s ovim generaliziranjem htio dobiti širi okvir za istraživanje, no rezultati bi možda bili jasniji za interpretiranje da je istraživanje orijentirano na jedan konkretan alat kako je to slučaj u nekim drugim istraživanjima (Biloš i Budimir, 2024; Strzelecki, 2023; Sobaih i sur., 2024; Sinanga i sur., 2024; Lavidas i sur., 2024; Lai i sur., 2024; Foroughi i sur., 2024).

U konačnici, treba uzeti u obzir kako je istraživačko područje prihvaćanja i korištenja generativne umjetne inteligencije relativno novo te da je još uvijek u formativnoj fazi zbog čega zahtijeva kontinuirano istraživanje ove teme. Prikupljeni podatci u ovoj relativno ranoj fazi dostupnosti i implementacije generativne umjetne inteligencije možda su samo odraz privremene percepcije koja se s vremenom sve veće prisutnosti, ali i daljnjom evolucijom ove tehnologije, može mijenjati i to predstavlja vremensko ograničenje ovog istraživanja, ali i svih srodnih studija.

8.6. Preporuke za buduća istraživanja

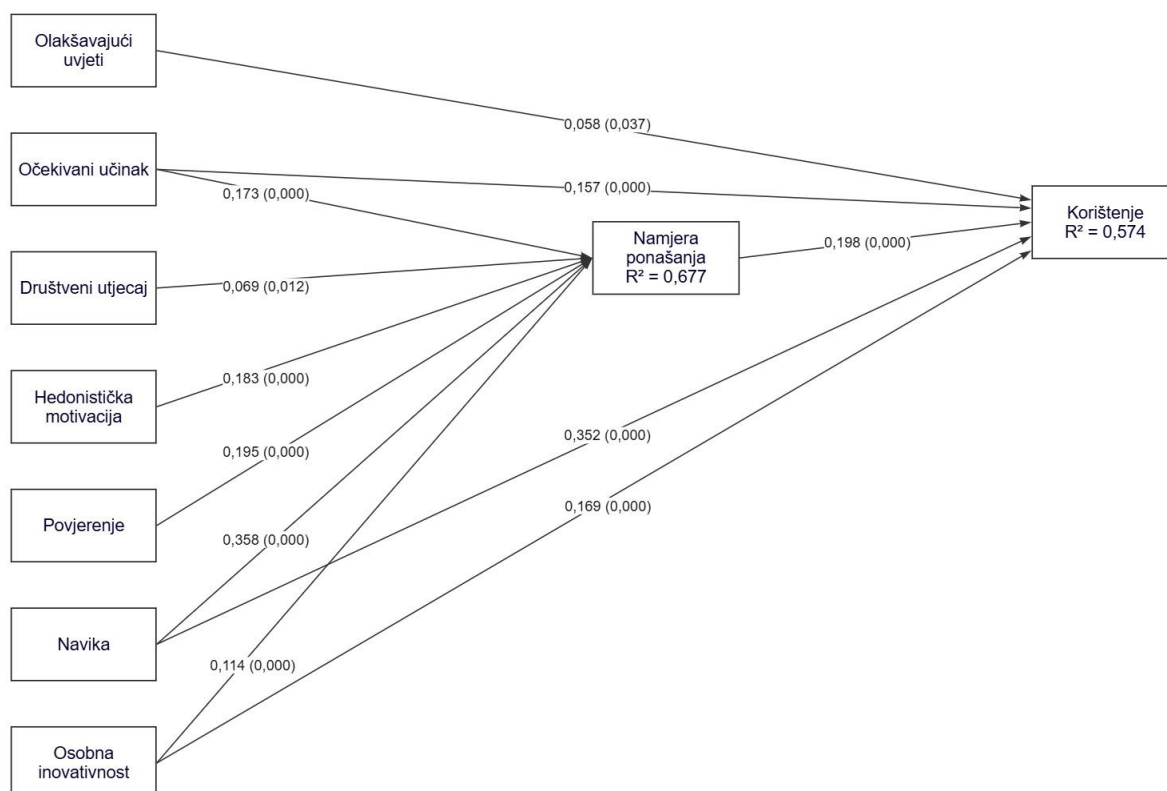
Temeljem provedenog istraživanja, dobivenih rezultata, ali i prepoznatih ograničenja, može se formulirati niz konkretnih preporuka za buduća istraživanja. Provođenje longitudinalnih istraživanja i panel studija koji bi pratili dinamiku prihvaćanja i korištenja generativne umjetne inteligencije svakako je jedan od glavnih prijedloga svim istraživačima unutar ovog

istraživačkog područja. Također, dubinska kvalitativna istraživanja mogla bi pružiti bogatije i kontekstualno opširnije uvide u motivaciju korisnika, ali i njihovo iskustvo korištenja.

Nadalje, buduća istraživanja mogla bi se usredotočiti na analizu prihvatanja specifičnih alata generativne umjetne inteligencije ili specifičnih područja primjene. Važno je i razmotriti te empirijski testirati uključivanje novih konstrukta, poput percipiranih etičkih rizika, anksioznosti u vezi korištenja umjetne inteligencije ili pak novih moderatora poput razine digitalne pismenosti. S obzirom na vrlo mali broj tradicionalnih moderatora koji su u okviru ovog istraživanja potvrđeni kao značajni, svakako je preporuka istraživanja novih potencijalnih moderatora koji bi mogli značajno unaprijediti prediktivnu snagu modela. Istraživanja na demografski raznolikijim uzorcima i različitim kulturološkim kontekstima također mogu proširiti postojeće znanstvene spoznaje na ovu temu, a tu se posebno može staviti naglasak na istraživanje prihvatanja i korištenja generativne umjetne inteligencije u nerazvijenim zemljama koje nisu obuhvaćene ovim istraživanjem. U konačnici, preporuka za istraživače iz drugih znanstvenih polja, koji bi mogli testirati utjecaj alata generativne umjetne inteligencije na kognitivne sposobnosti, kreativnost, navike pojedinaca, tržište rada ili efikasnost poslovnih subjekata mogle bi biti kompatibilne ovakvim istraživanjima jer nadopunju kontekst i važnost ovog istraživačkog polja.

Rezultatima empirijskog istraživanja potvrđeno je kako očekivani naponi (EE), cijena (PV) i olakšavajući uvjeti (FC) nisu statistički značajni prediktori namjere ponašanja (BI), ali i da moderatori u većoj mjeri nisu značajno utjecali na prediktivnu snagu konceptualnog modela. Shodno takvim rezultatima, autor je kreirao skraćeni model (slika 23.) bez moderirajućih efekata koji je značajno jednostavniji od postojećeg konceptualnog modela, ali mu je prediktivna snaga približno jednaka izvornom UTAUT2 modelu u kontekstu istraživanja prihvatanja i korištenja generativne umjetne inteligencije. Odnosno, objašnjenje varijance za BI je 1,2 postotna poena veća, dok je kod USE manja za 0,4 postotna poena u odnosu na UTAUT2 kakvog navode Venkatesh i sur. (2012). Zanimljivo je za uočiti kako su svi odnosi u takvom modelu statistički značajni te kako je *model fit* SRMR zadovoljavajući, 0,057 za cjeloviti, ali i za specificirani model.

Kako je vidljivo na slici 23., iz modela su izbačeni konstrukti EE i PV koji su se kontinuirano pokazivali neznačajnima za istraživanje u kontekstu generativne umjetne inteligencije. Jednako tako, iz modela je izbačena je i pretpostavka da FC utječe na BI, ali je ostavljen utjecaj FC na USE. U konačnici, iz modela su izbačeni i moderatori čiji je utjecaj u modelu vrlo ograničen, a prediktivna snaga vrlo slaba.



Slika 23. Predloženi skraćeni model

Premda teorijski okvir kojeg su predstavili Venkatesh i sur. (2012) ne predviđa utjecaj očekivanog učinka (PE) na korištenje (USE), navedeni je model podložen modifikacijama u svrhu prilagođavanja kontekstu. Tako primjerice Sobaih i sur. (2024) te Sergeeva i sur. (2025) navode kako je očekivani učinak (PE) statistički značajan prediktor korištenja alata generativne umjetne inteligencije (USE). Isto tako, Sun i sur. (2023) navode kako osobna inovativnost (PI) pojedinca značajno utječe na stvarno korištenje (USE). Shodno tome, u skraćeni model kao preporuku za buduća istraživanja, dodana su spomenuta dva odnosa koja značajno podižu prediktivnu snagu modela.

Kako je prethodno spomenuto, u istraživačkom instrumentu ispitanici su pitani koriste li plaćenu verziju nekog alata generativne umjetne inteligencije na što je 23 % ispitanika dalo potvrdni odgovor. Kao preporuka budućeg istraživanja može biti i stavljanje naglaska na ispitanike koji koriste plaćene ili *pro* verzije spomenutih alata. Navedene *pro* verzije u pravilu nude dodatne funkcionalnosti, više dozvoljenih aktivnosti u danu te veću i bržu obradu podataka. Za potrebe ovog doktorskog rada, odnosno njegovih preporuka za buduće istraživanje, ova dihotomna varijabla dobit će naziv „pro“. Za potrebe obrade prikupljenih podataka vrijednost „Ne“ kodirana je kao 1, dok je vrijednost „Da“ kodirana kao 2. Ukoliko bi

se varijabla *pro* koristila kao moderator odnosa između navike i stvarnog korištenja, tada bi objašnjenje varijance za korištenje porasla s 57,7 % na 60,2 %. Na poduzorku iz Ujedinjenog Kraljevstva objašnjenje varijance raste s 62,9 % na čak 67 %, dok na hrvatskom poduzorku objašnjenje varijance za korištenje raste s 54,5 % na 60,2 %. Ipak, znakovito je za uočiti kako $Pro \times HT > USE$ nije statistički značajan odnos premda snažno pomaže prediktivnoj snazi modela. Nadalje, ako bi se prethodno navedenom skraćenom modelu dodao moderator *pro* uz moderirajući efekt na utjecaj $HT > BI$, ponovno bi se potvrdilo kako $Pro \times HT > USE$ nije statistički značajan, ali se objašnjenje varijance za korištenje (USE) značajno podiže s 57,4 % na 60,0 %. Ovaj nalaz može se protumačiti efektom nepovratnog troška (engl. *sunk cost effect*) (Bokulić i Bovan, 2011), koji kaže kako su pojedinci skloniji koristiti proizvode ili usluge koje su već platili, kao primjerice skijanje po lošijim vremenskim uvjetima, ako je skijaška karta već unaprijed plaćena ili dodatnim teniskim lekcijama unatoč boli od ozljede (Soman i Cheema, 2001; Soman i Gourville, 2001). Ipak, navedeni odnos može biti objašnjen i teorijom samoodređenja jer plaćanje *pro* verzije alata može povećati autonomiju i osjećaj kontrole, što može rezultirati da korisnik češće i intenzivnije koristi alat, ne samo iz navike nego i iz percipirane koristi i vrijednosti (Ryan i Deci, 2000). Istraživanje ovog fenomena svakako je preporuka za sve istraživače koji će se baviti ovom i povezanim temama.

U konačnici, iz empirijskih je nalaza vidljivo kako očekivani napor (EE) nije važan prediktor namjere ponašanja (BI) i teško je za očekivati da će alati postati teži za korištenje, veća je vjerojatno da će postati sve jednostavniji i intuitivniji, a korisnici sve naviknutiji na njih. Sukladno navedenom, autor predlaže izbacivanje konstrukta EE iz budućih istraživanja prihvaćanja i korištenja generativne umjetne inteligencije mjerene UTAUT2 modelom. Cijena je faktor koji bi potencijalno ponovno mogao postati važan faktor ukoliko pružatelji ovih digitalnih usluga dodatno ograniče mogućnosti besplatnih verzija i time potaknu korisnike na plaćanje alata. Preporuka za istraživače je i da se uključuje dodatni relevantni faktori jer istraživanja potvrđuju kako faktori poput anksioznosti (Budhathoki i sur., 2024), straha od tehnologije, (Leschanowsky i sur., 2024), održivosti (Gansser & Reich, 2021), percipirane privatnosti (Merhi i sur., 2019), kvalitete sadržaja (Tantra & Ariyanti, 2017) ili znatiželje (Sinaga i sur., 2024) mogu imati značajnu ulogu u predikciji namjere ponašanja i stvarnog korištenja generativne umjetne inteligencije i srodnih tehnologija.

8.7. Osvrt na korištenje generativne umjetne inteligencije u pisanju doktorskog rada

U procesu pisanja ovog doktorskog rada korišteni su alati generativne umjetne inteligencije u skladu sa strogim pravilima najprestižnijih visokoobrazovnih institucija te akademskih nakladnika poput Springera, Springer Naturea, Wileya, SAGE-a, Elseviera, MDPI-ja, Taylor & Francisa te po etičkim uputama Europske unije o odgovornom korištenju generativne umjetne inteligencije u istraživanjima.

Istraživači sa Sveučilišta u Oxfordu, Sveučilišta u Cambridgeu, Sveučilišta u Kopenhagenu, Nacionalnog sveučilišta u Singapuru i drugih vodećih institucija osmislili su filozofski utemeljene etičke smjernice za korištenje velikih jezičnih modela u akademskom pisanju. Navode važnost definiranja korištenih alata generativne umjetne inteligencije koji će omogućiti veću transparentnost, ali i pojednostaviti pridržavanje etičkih standarda u akademskom pisanju. Tri su ključna kriterija kojih se autori trebaju pridržavati (Porsdam Mann i sur., 2024; University of Oxford, 2024):

- ljudska provjera te jamstvo točnosti i integriteta
- osiguravanje značajnog ljudskog doprinosa radu
- odgovarajuće priznanje i transparentnost korištenja generativne umjetne inteligencije.

Na sličnom tragu su i smjernice Europske komisije (2025) koje navode kako je istraživač u potpunosti odgovoran za integritet sadržaja kreiranog uz pomoć generativne umjetne inteligencije. Alati generativne umjetne inteligencije nerijetko su pristrani, haluciniraju i daju netočne podatke zbog čega je kritički pristup i provjera svega nužna. Spomenuti alati generativne umjetne inteligencije ne mogu i ne smiju biti autori niti koautori znanstvenih i stručnih radova zbog čega je odgovornost za napisani rad potpuno na ljudima. Smjericama Europske unije potiče se istraživače da transparentno navode za što su sve korišteni alati generativne umjetne inteligencije u procesu pisanja. Istraživači trebaju biti svjesni da sadržaj koji stavljaju u alate generativne umjetne inteligencije može biti korišten za trening AI modela što je vrlo opasno ukoliko se u alat stavlja sadržaj vrlo osobne prirode ili određene poslovne tajne.

Nalik navedenim preporukama, vodeće nakladničke kuće kreirale su svoje upute za autore. SAGE (2025) navodi kako alati generativne umjetne inteligencije smiju biti korišteni za unaprjeđenje čitljivosti napisanog sadržaja, prevođenje teksta te obradu kvantitativnih i

kvalitativnih podataka pod uvjetom da autori provjere točnost svega napisanog i sve transparentno prikažu. Na istom tragu je i Springer (2025) koji navodi kako je uređivanje teksta potpomognuto umjetnom inteligencijom s ciljem poboljšanja stila i čitljivost te smanjenja gramatičkih i pravopisnih grešaka, potpuno legitimno, ali da je odgovornost za izvorni rad potpuno na autorima. Springer Nature (2025) dodatno upozorava na korištenje inkluzivnih skupova podataka s ciljem izbjegavanja algoritamske pristranosti jer u suprotnom bi se ovaj štetni utjecaj alata umjetne inteligencije mogao pojačati ukoliko model bude treniran na tim istim problematičnim podacima. Taylor & Francis (2025) navode kako autori moraju jasno naznačiti puni naziv korištenih alata (s brojem verzije), kako je korišten i s kojim razlogom te da bilo koje korištenje generativne umjetne inteligencije treba biti provedeno uz ljudski nadzor, etično i transparentno. Elsevier (2025) navodi kako autori smiju koristiti alate generativne umjetne inteligencije za poboljšanje čitljivosti rada, ali ne za izvođenje znanstvenih zaključaka. Sav sadržaj u radu odgovornost je autora, a za bilo kakvo korištenje alata generativne umjetne inteligencije autori trebaju napisati izjavu o korištenim alatima. Takva izjava objavljuje se u radu te podržava transparentnost i povjerenje između autora, čitatelja, recenzenata i urednika. Slične standarde uveli su i Wiley (2025) te MDPI (2025).

Sukladno prethodno navedenim smjernicama vodećih akademskih institucija, s ciljem transparentnosti, autor navodi kako su u pisanju ovog doktorskog rada korišteni ChatGPT (GPT-3.5, GPT-4o i GPT-4.5) kod prevođenja, sažimanja, stilskog obogaćivanja i uređivanja tekstova te za konzultiranje o statističkim i metodološkim procedurama. Za istu svrhu korišteni su i Gemini (2.0 Flash) te Google AI studio (Gemini 2.5 Pro Preview). Za pronalazak stručne literature, informacija i izvješća korišten je Perplexity Pro. S druge strane, za pomoć u pronalasku znanstvene literature korišten je SciSpace. Nadalje, pripremu podataka prikupljenih primarnim istraživanjem i kodiranje vrijednosti unutar te iste baze podataka korišten je Data Analyst, GPT ekstenzija unutar ChatGPT-a. Alati generativne umjetne inteligencije korišteni su etično i odgovorno. Sav kreirani sadržaj potpuna je odgovornost autora doktorskog rada.

9. ZAKLJUČAK

Ovaj doktorski rad pruža sveobuhvatni uvid u odrednice prihvaćanja i korištenja generativne umjetne inteligencije mjeren proširenim UTAUT2 modelom na ispitanicima iz Ujedinjenog Kraljevstva i Republike Hrvatske. U eri koja se naziva i „AI proljeće“, a koju obilježava eksponencijalni rast i razvoj te sveprisutna integracija umjetne inteligencije na mnogobrojnim poljima privatnog i poslovnog života, generativna umjetna inteligencija izdvaja se kao tehnologija s potencijalom da snažno transformira način na koji pojedinci rade, uče i komuniciraju međusobno, ali i s drugim tehnologijama. Upravo iz navedenog razloga ovaj doktorski rad predstavlja značajan i pravovremen doprinos u istraživanju odrednica prihvaćanja i korištenja generativne umjetne inteligencije. Potencijal ove tehnologije i njezin socio-ekonomski utjecaj snažno je povezan s prihvaćanjem i korištenjem od strane krajnjih korisnika i upravo je zato važno razviti što precizniji model mjerenja namjere korištenja i stvarnog korištenja generativne umjetne inteligencije, ali i kontinuirano pratiti dinamiku korisničkih navika. U okviru ovog doktorskog rada temeljito je analiziran, a potom i proširen UTAUT2 model kao mjerni instrument s ciljem prilagodbe konceptualnog modela za istraživanja u kontekstu generativne umjetne inteligencije. U konačnici, prilagođeni je mjerni instrument korišten kao teorijski okvir na kojem je provedeno empirijsko istraživanje na komparativnim uzorcima iz Ujedinjenog Kraljevstva i Republike Hrvatske.

Osnovni istraživački problem kojim se bavio ovaj doktorski rad jest identificiranje i kvantificiranje ključnih odrednica u prihvaćanju i korištenju generativne umjetne inteligencije. Dok alati generativne umjetne inteligencije, poput ChatGPT-a ili Geminija bilježe snažan rast popularnosti, ali i primjene, znanstvena zajednica još uvijek traži idealan konceptualni model za mjerenje odrednica u prihvaćanju i korištenju takvih alata. Generativna umjetna inteligencija pred opću javnost, ali i akademske istraživače stavlja nove izazove jer je to prva tehnologija koja ima mogućnost generiranja sadržaja koji je toliko sofisticiran da ga je ponekad gotovo nemoguće razlikovati od sadržaja koji kreira sam čovjek. Naravno, takvi tehnološki iskoraci sa sobom nose razne etičke dileme, ali i zabrinutost javnosti oko pouzdanosti, privatnosti, autorskih prava te mnogih drugih elemenata koji bi mogli imati negativan utjecaj na društvo.

Modeli prihvaćanja tehnologije, uključujući i UTAUT2, pružaju čvrstu teorijsku osnovu o glavnim odrednicama koje utječu na prihvaćanje tehnologije. Ipak, takve modele često treba prilagođavati kontekstu istraživanja s ciljem boljeg razumijevanja čimbenika koji utječu na namjeru korištenja i korištenje određene tehnologije. Shodno tomu, glavni cilj ovog doktorskog

rada bio je razviti, empirijski testirati i validirati novonastali konceptualni model, odnosno prošireni UTAUT2 model. Kroz njega, cilj je bio identificirati ključne odrednice koje utječu na namjeru ponašanja i stvarno ponašanje ovih naprednih alata. Osim toga, jedna od važnijih karakteristika ovog rada je usporedna analiza na dva različita socio-ekonomska, tehnološka i kulturna konteksta - Hrvatska kao tržište u razvoju po pitanju spremnosti na umjetnu inteligenciju i Ujedinjeno Kraljevstvo kao jedno od tehnološki najrazvijenijih zemalja svijeta koja je izrazito spremna na promjene koje sa sobom donosi umjetna inteligencija.

U slučaju ovog doktorskog rada gdje je naglasak stavljen na generativnu umjetnu inteligenciju, konceptualni model kreiran je nakon temeljite analize postojeće literature te proširen konstruktima osobne inovativnosti i povjerenja. Cilj ovog doktorskog rada bio je analizirati razvoj modela UTAUT2 i njegovu integraciju u različitim kontekstima, na različitim tržištima i u različitim formama te shodno tim zaključcima predložiti novi, kontekstualni prilagođeni model, a potom ga i empirijski testirati te ponuditi teorijske i aplikativne doprinose te zaključke takvog istraživanja.

Na temelju empirijskih rezultata može se zaključiti kako konceptualni model ovog rada, u formi proširenog i prilagođenog UTAUT2 modela, objašnjava 70,1 % varijance u namjeri ponašanja (BI) i 57,7 % varijance u stvarnom korištenju (USE) pri obradi cjelovitog uzorka. Ključni prediktori namjere ponašanja, koji se kontinuirano potvrđuju kao značajni, su navika (HT), hedonistička motivacija (HM), povjerenje (TR), očekivani učinak (PE), društveni utjecaj (SI) i osobna inovativnost (PI). S druge strane, očekivani naponi (EE) i cijena (PV) uglavnom nisu značajni prediktori namjere ponašanja, dok olakšavajući uvjeti (FC) nisu utjecali na namjeru ponašanja (BI), ali jesu na stvarno korištenje (USE). Utjecaj na stvarno korištenje potvrđen je i za navike (HT) te namjeru ponašanja (BI). Beznačajnost očekivanih napora (EE) sugerira kako korisnici alate generativne umjetne inteligencije doživljavaju jednostavnima za korištenje za koje nije potrebno uložiti dodatan napor. S druge strane, cijenu (PV) ne doživljaju značajnim prediktorom namjere korištenja jer su alati generativne umjetne inteligencije u većoj mjeri dostupni i u besplatnoj verziji što sugerira da utjecaj monetarnog troška ne predstavlja prepreku u odluci o korištenju takvih alata. Uloga moderatora, dobi, spola i iskustva, pokazala se vrlo ograničenom što sugerira kako demografske karakteristike i prethodno iskustvo nemaju značajnih utjecaja na odluku o korištenju alata generativne umjetne inteligencije.

Također, analiza po tržištima pokazala je važne kulturološke razlike pri čemu su ispitanici iz Ujedinjenog Kraljevstva, kao tehnološki naprednijeg društva, više vođeni koristima generativne umjetne inteligencije, ali i vlastite proaktivnosti po pitanju korištenja tehnologije. Nadalje,

analiza je pokazala kako je u Ujedinjenom Kraljevstvu, pored univerzalnih intrinzičnih faktora poput zabave, rutine i povjerenja, snažan i utjecaj društvene okoline. S druge strane, u Hrvatskoj se prediktori očekivani učinak (PE), društveni utjecaj (SI) i osobna inovativnost (PI) nisu potvrdili kao značajni. U Hrvatskoj su dominantni prediktori namjere ponašanja navika (HT), hedonistička motivacija (HM) i povjerenje (TR). Zaključak je kako su hrvatski korisnici u fazi formiranja namjere korištenja generativne umjetne inteligencije primarno vođeni unutarnjim motivima poput užitka, navike i fundamentalnog osjećaja sigurnosti kod korištenja tehnologije, dok su vanjski faktori poput percipirane korisnosti i društvenog utjecaja bili slabiji ili uvjetovati demografskim specifičnostima, a što na neki način odražava zašto je Hrvatska još uvijek u društvu zemalja u razvoju po pitanju spremnosti na umjetnu inteligenciju.

Provedeno empirijsko istraživanje generira nekoliko važnih teorijskih implikacija za istraživanja u domeni prihvaćanja tehnologije, s naglaskom na nove i napredne tehnologije poput generativne umjetne inteligencije. Rad uspješno proširuje i prilagođava UTAUT2 model za kontekst generativne umjetne inteligencije kroz integraciju konstrukta osobne inovativnosti i povjerenja, no ujedno i dokazuje fleksibilnost UTAUT2 modela kao teorijskog okvira. Osim glavnog konceptualnog modela, u preporukama istraživanja stavljeni su i prijedlozi kako model reducirati te dodatno prilagoditi spomenutom kontekstu.

U konačnici, ovo empirijsko istraživanje, kao i gotovo svako opsežnije i kompleksnije istraživanje, ima svoja ograničenja, no njegovi nalazi nude značajne teorijske i aplikativne doprinose. Ove važne spoznaje imaju implikacije na širok spektar dionika, od marketinških stručnjaka, obrazovnih djelatnika, tvoraca politika pa sve do samih korisnika. Predloženi su smjerovi za daljnja istraživanja koja mogu osigurati kontinuitet, ali i daljnje produbljivanje znanstvenih spoznaja u ovom izrazito dinamičnom i društveno relevantnom području. Generativna umjetna inteligencija zasigurno će nastaviti transformirati privatno i poslovno okruženje pojedinaca, a znanstveno utemeljeno razumijevanje razloga zbog kojih ljudi koriste ovu tehnologiju, ali i načina na koji koriste generativnu umjetnu inteligenciju ostat će od presudne važnosti za usmjeravanje njezinog razvoja, ali i ostvarivanje njenog punog potencijala za dobrobit društva. Ovaj doktorski rad predstavlja važan korak u boljem razumijevanju o čimbenicima koji utječu na prihvaćanje i korištenje generativne umjetne inteligencije, ali i uspoređivanje korisnika iz vlastitog okruženja s korisnicima tehnološki naprednije ekonomije kakva težimo postati.

Zaključno, na temelju pomno proučene relevantne literature i provedenog empirijskog istraživanja, autor rada vjeruje kako disruptivna tehnologija u središtu ovog istraživanja iz

temelja mijenja privatne i poslovne aktivnosti pojedinaca diljem svijeta. Iako je relativno nova, koristi ju prilično veliki broj korisnika koji u simbiozi s njom ostvaruju značajno veću učinkovitost nego oni koji još nisu prigrlili ovu novu tehnologiju. Integracija generativne umjetne inteligencije uskoro će biti vidljiva u mnogobrojnim digitalnim platformama i postat će sastavna funkcionalnost gotovo svega u internetskom okruženju. Strah od negativnih utjecaja je opravdan, posebice jer je riječ o možda najmoćnijoj tehnologiji ikada. Ipak, ovakvim i komplementarnim istraživanjima znanost može dati svoj obol zajednici. Nužno je kontinuirano pratiti razvoj umjetne inteligencije, ali i korisničko ponašanje kod upotrebe iste. Također, nužno je pomoći svim tržišnim segmentima da prigrle generativnu umjetnu inteligenciju te srodne tehnologije, pomoći tvorcima politika da stvore sigurno okruženje i osiguraju etičnu uporabu tehnologije kako bismo se kao društvo pripremili za *vrli novi svijet*, svijet u kojem će se sutra možda raditi manje ako društvo bude dovoljno spretno u rukovanju umjetnom inteligencijom.

10. POPIS LITERATURE

1. Agarwal, R., & Prasad, J. (1997). The role of innovation characteristics and perceived voluntariness in the acceptance of information technologies. *Decision sciences*, 28 (3), 557–582.
2. Agarwal, R.; Prasad, J. (1998) A conceptual and operational definition of personal innovativeness in the domain of information technology. *Inf. Syst. Res.*, 9, 204–215.
3. Aggarwal, A., Mittal, M., & Battineni, G. (2021). Generative adversarial network: An overview of theory and applications. *International Journal of Information Management Data Insights*, 1 (1), 100004.
4. Ain, N., Kaur, K., & Waheed, M. (2016). The influence of learning value on learning management system use: An extension of UTAUT2. *Information Development*, 32 (5), 1306–1321.
5. Ajzen, I. (1991). The theory of planned behavior. *Organizational behavior and human decision processes*, 50 (2), 179–211.
6. Alba, J. W., & Hutchinson, J. W. (1987). Dimensions of consumer expertise. *Journal of consumer research*, 13 (4), 411–454.
7. Aldaroso, I., Armantier, O., Doerr, S., Gambacorta, L. & Oliviero, T. (2024a). The gen AI gender gap. BIS Working Paper, No. 1197. Dostupno na: <https://www.bis.org/publ/work1197.pdf> [Pristupljeno: 3. 3. 2025.]
8. Aldaroso, I., Armantier, O., Doerr, S., Gambacorta, L. & Oliviero, T. (2024b). *Survey evidence on gen AI and households: job prospects amid trust concerns*. BIS Bulletin, No 86. Dostupno na: <https://www.bis.org/publ/bisbull86.pdf> [Pristupljeno: 3. 3. 2025.]
9. Al-Emran, M., AlQudah, A. A., Abbasi, G. A., Al-Sharafi, M. A., & Iranmanesh, M. (2023). Determinants of using AI-based chatbots for knowledge sharing: evidence from PLS-SEM and fuzzy sets (fsQCA). *IEEE Transactions on Engineering Management*.
10. Ali, M. S. M., Wasel, K. Z. A., & Abdelhamid, A. M. M. (2024). Generative AI and Media Content Creation: Investigating the Factors Shaping User Acceptance in the Arab Gulf States. *Journalism and Media*, 5 (4), 1624–1645.

11. Al-kfairy, M., Mustafa, D., Kshetri, N., Insiew, M., & Alfandi, O. (2024, September). Ethical challenges and solutions of generative AI: An interdisciplinary perspective. In *Informatics* (Vol. 11, No. 3, p. 58). Multidisciplinary Digital Publishing Institute.
12. Ameri, A., Khajouei, R., Ameri, A., & Jahani, Y. (2020). Acceptance of a mobile-based educational application (LabSafety) by pharmacy students: An application of the UTAUT2 model. *Education and Information Technologies*, 25 (1), 419–435.
13. Baabdullah, A. M. (2018). Consumer adoption of Mobile Social Network Games (M-SNGs) in Saudi Arabia: The role of social influence, hedonic motivation and trust. *Technology in society*, 53, 91–102.
14. Bakan, D. (1966). The duality of human existence: An essay on psychology and religion.
15. Bandura, A. (1986). Social foundations of thought and action. *Englewood Cliffs, NJ*, 1986, 23–28.
16. Banh, L., & Strobel, G. (2023). Generative artificial intelligence. *Electronic Markets*, 33 (1), 63.
17. Baptista, G., & Oliveira, T. (2015). Understanding mobile banking: The unified theory of acceptance and use of technology combined with cultural moderators. *Computers in human behavior*, 50, 418–430.
18. Barclay, D. W., Higgins, C. A., & Thompson, R. (1995). The partial least squares approach to causal modeling: Personal computer adoption and use as illustration. *Technology Studies*, 2, 285–309.
19. Baudier, P., Ammi, C., & Deboeuf-Rouchon, M. (2020). Smart home: Highly-educated students' acceptance. *Technological Forecasting and Social Change*, 153, 119355.
20. Bellet, T., & Banet, A. (2023). UTAUT4-AV: An extension of the UTAUT model to study intention to use automated shuttles and the societal acceptance of different types of automated vehicles. Transportation research part F: *Traffic psychology and behaviour*, 99, 239–261.
21. Bem, D. J., & Allen, A. (1974). On predicting some of the people some of the time: The search for cross-situational consistencies in behavior. *Psychological review*, 81 (6), 506.
22. Bick, A.; Blandin, A. & Deming, D.J. (2024). *The Rapid Adoption of Generative AI*. National Bureau of Economic Research. Dostupno na:

- https://www.nber.org/system/files/working_papers/w32966/w32966.pdf [Pristupljeno: 19. 2. 2024.]
23. Biloš, A., & Budimir, B. (2024). Understanding the Adoption Dynamics of ChatGPT among Generation Z: Insights from a Modified UTAUT2 Model. *Journal of theoretical and applied electronic commerce research*, 19 (2), 863–879.
24. Bin-Nashwan, S. A., Sadallah, M., & Bouteraa, M. (2023). Use of ChatGPT in academia: Academic integrity hangs in the balance. *Technology in Society*, 75, 102370.
25. Blackman, R. (2020). *A Practical Guide to Building Ethical AI*. Harvard Business Review. Dostupno na: <https://hbr.org/2020/10/a-practical-guide-to-building-ethical-ai> [Pristupljeno: 30. 11. 2024.]
26. Bloomberg (2023). *Generative AI to Become a \$1,3 Trillion Market by 2032*, Research Finds. Bloomberg. Dostupno na: <https://www.bloomberg.com/company/press/generative-ai-to-become-a-1-3-trillion-market-by-2032-research-finds/> [Pristupljeno: 17. 7. 2024.]
27. Blut, M., Wang, C., Wunderlich, N. V., & Brock, C. (2021). Understanding anthropomorphism in service provision: a meta-analysis of physical robots, chatbots, and other AI. *Journal of the Academy of Marketing Science*, 49, 632–658.
28. Bommasani, R. (2023). *AI Spring? Four Takeaways from Major Releases in Foundation Models*. Stanford University Human-Centered Artificial Intelligence. Dostupno na: <https://hai.stanford.edu/news/ai-spring-four-takeaways-major-releases-foundation-models> [Pristupljeno: 3. 12. 2024.]
29. Bostrom, N. (2014). *Superintelligence: Paths, dangers, strategies*. United Kingdom. Oxford University Press.
30. Bozionelos, N. (1996). Psychology of computer use: XXXIX. Prevalence of computer anxiety in British managers and professionals. *Psychological reports*, 78(3), 995-1002.
31. Broady, T., Chan, A., & Caputi, P. (2010). Comparison of older and younger adults' attitudes towards and abilities with computers: Implications for training and learning. *British Journal of Educational Technology*, 41 (3), 473–485.
32. Brown, S. A., & Venkatesh, V. (2005). Model of adoption of technology in households: A baseline model test and extension incorporating household life cycle. *MIS quarterly*, 399–426.

33. Brown, S. A., & Venkatesh, V. (2005). Model of adoption of technology in households: A baseline model test and extension incorporating household life cycle. *MIS quarterly*, 399–426.
34. Browne, R. (2024). *Britain looks to upstage France with play for world's third major AI hub after U.S., China*. CNBC. Dostupno na: <https://www.cnn.com/2024/06/26/britain-looks-to-upstage-france-with-play-for-worlds-third-major-ai-hub-.html> [Pristupljeno: 30. 10. 2024.]
35. Brühl, V. (2024). Generative Artificial Intelligence—Foundations, Use Cases and Economic Potential. *Intereconomics*, 59 (1), 5–9.
36. Bruschi, I., & Rappell, N. (2020). Exploring the acceptance of instant shopping—An empirical analysis of the determinants of user intention. *Journal of Retailing and Consumer Services*, 54, 101936.
37. Budhathoki, T., Zirar, A., Njoya, E. T., & Timsina, A. (2024). ChatGPT adoption and anxiety: a cross-country analysis utilising the unified theory of acceptance and use of technology (UTAUT). *Studies in Higher Education*, 1–16.
38. Cabero-Almenara, J., Palacios-Rodríguez, A., Rojas Guzmán, H. D. L. Á., & Fernández-Scagliusi, V. (2025). Prediction of the Use of Generative Artificial Intelligence Through ChatGPT Among Costa Rican University Students: A PLS Model Based on UTAUT2. *Applied Sciences*, 15 (6), 3363.
39. Cabrera-Sánchez, J. P., Villarejo-Ramos, Á. F., Liébana-Cabanillas, F., & Shaikh, A. A. (2021). Identifying relevant segments of AI applications adopters—Expanding the UTAUT2's variables. *Telematics and Informatics*, 58, 101529.
40. Cambridge Dictionary (n.d.). *Intelligence*. Cambridge Business English Dictionary. Cambridge University Press. Dostupno na: <https://dictionary.cambridge.org/dictionary/english/intelligence> [Pristupljeno: 28. 8. 2024.]
41. Cao, Y., Aziz, A. A., & Arshad, W. N. R. M. (2023). University students' perspectives on Artificial Intelligence: A survey of attitudes and awareness among Interior Architecture students. *IJERI: International Journal of Educational Research and Innovation*, (20), 28–49.
42. Chacko, H. E., Williams, K., & Schaffer, J. (2012). A conceptual framework for attracting Generation Y to the hotel industry using a seamless hotel organizational structure. *Journal of Human Resources in Hospitality & Tourism*, 11 (2), 106–122.

43. Chan, K. Y., Gong, M., Xu, Y., & Thong, J. (2008). Examining user acceptance of SMS: An empirical study in China and Hong Kong. *PACIS 2008 Proceedings*, 294.
44. Chatterjee, S., & Bhattacharjee, K. K. (2020). Adoption of artificial intelligence in higher education: A quantitative analysis using structural equation modelling. *Education and Information Technologies*, 25, 3443–3463.
45. Chau, P. Y., & Hui, K. L. (1998). Identifying early adopters of new IT products: A case of Windows 95. *Information & management*, 33 (5), 225–230.
46. Chin, W. W. (1998). "The Partial Least Squares Approach to Structural Equation Modeling." In *Modern Methods for Business Research* (ed. G. A. Marcoulides), Lawrence Erlbaum Associates.
47. Chiu, T. K. (2024). The impact of Generative AI (GenAI) on practices, policies and research direction in education: A case of ChatGPT and Midjourney. *Interactive Learning Environments*, 32 (10), 6187–6203.
48. Chopdar, P. K., Korfiatis, N., Sivakumar, V. J., & Lytras, M. D. (2018). Mobile shopping apps adoption and perceived risks: A cross-country perspective utilizing the Unified Theory of Acceptance and Use of Technology. *Computers in Human Behavior*, 86, 109–128.
49. Choudhary, S., Kaushik, N., Sivathanu, B., & Rana, N. P. (2024). Assessing Factors Influencing Customers' Adoption of AI-Based Voice Assistants. *Journal of Computer Information Systems*, 1–18.
50. Choung, H., David, P., & Ross, A. (2023). Trust and ethics in AI. *AI & Society*, 38 (2), 733–745.
51. Chow, A. & Perrigo, B. (2023). *The AI Arms Race Is Changing Everything*. Time. Dostupno na: <https://time.com/6255952/ai-impact-chatgpt-microsoft-google/> [Pristupljeno: 3. 12. 2024.]
52. 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.
53. Chu, T. H., Chao, C. M., Liu, H. H., & Chen, D. F. (2022). Developing an extended theory of UTAUT 2 model to explore factors influencing Taiwanese consumer adoption of intelligent elevators. *Sage Open*, 12 (4), 21582440221142209.

54. Cimperman, M., Brenčič, M. M., Trkman, P., & Stanonik, M. D. L. (2013). Older adults' perceptions of home telehealth services. *Telemedicine and e-Health*, 19 (10), 786–790.
55. Cintrón, J. J. V. (2022). *Factors Influencing IT Managers' Acceptance of Artificial Intelligence (AI) in Digital Transformation*. Capella University.
56. Cisco (2023). *Cisco AI Readiness Index*. Cisco. Dostupno na: https://www.cisco.com/c/dam/m/en_us/solutions/ai/readiness-index/documents/cisco-global-ai-readiness-index.pdf [Pristupljeno: 27. 10. 2024.]
57. Compeau, D. R., & Higgins, C. A. (1995). Computer self-efficacy: Development of a measure and initial test. *MIS quarterly*, 189–211.
58. Corchado, J. M., López, S., Garcia, R., & Chamoso, P. (2023). Generative artificial intelligence: fundamentals. *ADCAIJ: advances in distributed computing and artificial intelligence journal*, 12 (1), e31704–e31704.
59. Corchado, J. M., López, S., Garcia, R., & Chamoso, P. (2023). Generative artificial intelligence: fundamentals. *ADCAIJ: advances in distributed computing and artificial intelligence journal*, 12 (1), e31704–e31704.
60. Cortez, P. M., Ong, A. K. S., Diaz, J. F. T., German, J. D., & Jagdeep, S. J. S. S. (2024). Analyzing Preceding factors affecting behavioral intention on communicational artificial intelligence as an educational tool. *Heliyon*, 10 (3).
61. Coulter, K. S., & Coulter, R. A. (2007). Distortion of price discount perceptions: The right digit effect. *Journal of Consumer Research*, 34 (2), 162-173.
62. Cronbach, L. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika* 16, 297–334. doi: 10.1007/BF02310555
63. D. Oliver (2007). The World According to Y: Inside the New Adult Generation by R. Huntley, *Labour History*, No. 92, pp. 186–188
64. Darley, W. K., & Smith, R. E. (1995). Gender differences in information processing strategies: An empirical test of the selectivity model in advertising response. *Journal of advertising*, 24 (1), 41–56.
65. Dash, G., & Paul, J. (2021). CB-SEM vs PLS-SEM methods for research in social sciences and technology forecasting. *Technological Forecasting and Social Change*, 173, 121092.
66. Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS quarterly*, 319–340.

67. Davis, F. D., & Venkatesh, V. (2004). Toward preprototype user acceptance testing of new information systems: implications for software project management. *IEEE Transactions on Engineering management*, 51 (1), 31–46.
68. Davis, F. D., Bagozzi, R. P., & Warshaw, P. R. (1989). User acceptance of computer technology: A comparison of two theoretical models. *Management science*, 35 (8), 982–1003.
69. Davis, F. D., Bagozzi, R. P., & Warshaw, P. R. (1992). Extrinsic and intrinsic motivation to use computers in the workplace 1. *Journal of applied social psychology*, 22 (14), 1111–1132.
70. Davison, A. C., and Hinkley, D. V. (1997). *Bootstrap Methods and Their Application*, Cambridge University Press: Cambridge.
71. Dawson, A. (2024). *AI could boost UK GDP by £550 billion by 2035, research shows*. Microsoft UK Stories. Dostupno na: <https://ukstories.microsoft.com/features/ai-could-boost-uk-gdp-by-550-billion-by-2035-research-shows/> [Pristupljeno: 7. 2. 2024.]
72. Deaux, K., & Lewis, L. L. (1984). Structure of gender stereotypes: Interrelationships among components and gender label. *Journal of personality and Social Psychology*, 46 (5), 991.
73. Deloitte (2024). *Now decides next: Insights from the leading edge of generative AI adoption*. Deloitte. Dostupno na: <https://www2.deloitte.com/content/dam/Deloitte/us/Documents/consulting/us-state-of-gen-ai-report.pdf> [Pristupljeno: 18. 9. 2024.]
74. Deloitte (2024). *Over 18 million people in the UK have now used Generative AI*. Deloitte. Dostupno na: <https://www.deloitte.com/uk/en/about/press-room/over-eighteen-million-people-in-the-uk-have-now-used-generative-ai.html> [Pristupljeno: 19. 2. 2024.]
75. Deloitte (2024). *Women and generative AI: The adoption gap is closing fast, but a trust gap persists*. Deloitte. Dostupno na: <https://www2.deloitte.com/us/en/insights/industry/technology/technology-media-and-telecom-predictions/2025/women-and-generative-ai.html#endnote-13> [Pristupljeno: 3. 3. 2025.]
76. Department for Science, Innovation and Technology UK (2025). *UK AI sector attracts £200 million a day in private investment since July*. Gov.uk. Dostupno na:

<https://www.gov.uk/government/news/uk-ai-sector-attracts-200-million-a-day-in-private-investment-since-july> [Pristupljeno: 7. 2. 2025.]

77. Dhariwal, P., & Nichol, A. (2021). Diffusion models beat gans on image synthesis. *Advances in neural information processing systems*, 34, 8780–8794.
78. Dijkstra, T. K. (2009). Latent variables and indices: Herman Wold's basic design and partial least squares. In *Handbook of partial least squares: Concepts, methods and applications* (pp. 23–46). Berlin, Heidelberg: Springer Berlin Heidelberg.
79. Dijkstra, T. K. (2014). PLS'Janus face—response to professor Rigdon's 'rethinking partial least squares Modeling: in praise of simple methods'. *Long Range Planning*, 47 (3), 146–153.
80. Dijkstra, T. K., & Henseler, J. (2015). Consistent partial least squares path modeling. *MIS quarterly*, 39 (2), 297–316.
81. Dimitriadis, C. (2024). *The Future of The Cybersecurity Profession With The Rise of AI*. Forbes. Dostupno na: <https://www.forbes.com/councils/forbestechcouncil/2024/07/03/the-future-of-the-cybersecurity-profession-with-the-rise-of-ai/> [Pristupljeno: 1. 12. 2024.]
82. Dinev, T., Albano, V., Xu, H., D'Atri, A., & Hart, P. (2016). Individuals' attitudes towards electronic health records: A privacy calculus perspective. *Advances in healthcare informatics and analytics*, 19–50.
83. Dissanayake, C. A. K., Jayathilake, W., Wickramasuriya, H. V. A., Dissanayake, U., Kopiyawattage, K. P. P., & Wasala, W. M. C. B. (2022). Theories and models of technology adoption in agricultural sector. *Human Behavior and Emerging Technologies*, 2022.
84. Dodds, W. B., Monroe, K. B., & Grewal, D. (1991). Effects of price, brand, and store information on buyers' product evaluations. *Journal of marketing research*, 28 (3), 307–319.
85. Duhigg, C. (2013). *The Power of Habit: Why we do what we do and how to change*. Random House.
86. DZS (2023). *Kontinuirani rast udjela visokoobrazovanog stanovništva*. Državni zavod za statistiku. Dostupno na: <https://dzs.gov.hr/vijesti/kontinuiran-rast-udjela-visokoobrazovanog-stanovnistva/1594> [Pristupljeno: 6. 5. 2025.]

87. Efron, B., and Tibshirani, R. J. (1993). *An Introduction to the Bootstrap*, Chapman Hall: New York.
88. Elias, S. M., Smith, W. L., & Barney, C. E. (2012). Age as a moderator of attitude towards technology in the workplace: Work motivation and overall job satisfaction. *Behaviour & Information Technology*, 31 (5), 453–467.
89. Emon, M. M. H., Hassan, F., Nahid, M. H., & Rattanawiboonsom, V. (2023). Predicting Adoption Intention of Artificial Intelligence-A Study on ChatGPT. *AIUB Journal of Science and Engineering (AJSE)*, 22 (2), 189–196.
90. Enciklopedija (n.d.). *Inteligencija*. Dostupno na: <https://enciklopedija.hr/natuknica.aspx?ID=27600> [Pristupljeno: 28. 8. 2024.]
91. EU Artificial Intelligence Act (2024). *Zakon EU o umjetnoj inteligenciji. EU Artificial Intelligence Act*. Dostupno na: <https://artificialintelligenceact.eu/> [Pristupljeno: 23. 10. 2024.]
92. Europska komisija (2018). *Komunikacija komisije Europskom parlamentu, Europskom vijeću, Vijeću, Europskom gospodarskom i socijalnom odboru i Odboru regija – Umjetna inteligencija za Europu*. Dokument 52018DC0237. Službena stranica Europske unije. Dostupno na: <https://eur-lex.europa.eu/legal-content/HR/TXT/?uri=COM:2018:237:FIN> [Pristupljeno: 30. 8. 2024.]
93. Europska komisija (2019). *Etičke smjernice za pouzdanu umjetnu inteligenciju*. Neovisna stručna skupina na visokoj razini o umjetnoj inteligenciji Europske komisije. Dostupno na: https://www.europarl.europa.eu/meetdocs/2014_2019/plmrep/COMMITTEES/JURI/DV/2019/11-06/Ethics-guidelines-AI_HR.pdf [7. 2. 2024.]
94. Europska komisija (2021). *Digitalno desetljeće Europe: digitalni ciljeva za 2030*. Europska komisija. Dostupno na: https://commission.europa.eu/strategy-and-policy/priorities-2019-2024/europe-fit-digital-age/europes-digital-decade-digital-targets-2030_hr [Pristupljeno: 23. 10. 2024.]
95. Europska komisija (2025). *Komisija objavila ambiciozan akcijski plan za Europu kao kontinent umjetne inteligencije*. Europska komisija. Dostupno na: https://croatia.representation.ec.europa.eu/news/komisija-objavila-ambiciozan-akcijski-plan-za-europu-kao-kontinent-umjetne-inteligencije-2025-04-09_hr [Pristupljeno: 17. 5. 2025.]

96. Europska komisija (2025). *Living guidelines on the responsible use of generative AI in research*. ERA Forum Stakeholders' document. European commission. Dostupno na: https://research-and-innovation.ec.europa.eu/document/download/2b6cf7e5-36ac-41cb-aab5-0d32050143dc_en?filename=ec_rtd_ai-guidelines.pdf [Pristupljeno: 13. 5. 2025.]
97. Europski parlament (2024a). *Artificial Intelligence Act: MEPs adopt landmark law*. European Parliament. Dostupno na: <https://www.europarl.europa.eu/news/en/press-room/20240308IPR19015/artificial-intelligence-act-meps-adopt-landmark-law> [Pristupljeno: 23. 10. 2024.]
98. Europski parlament (2024b). *Investment in artificial intelligence in the National Recovery and Resilience Plans*. Europski parlament. Dostupno na: [https://www.europarl.europa.eu/thinktank/en/document/EPRS_BRI\(2024\)762288](https://www.europarl.europa.eu/thinktank/en/document/EPRS_BRI(2024)762288) [Pristupljeno: 7. 2. 2024.]
99. Europski revizorski sud (2024). *Ambicije Europske unije o području umjetne inteligencije*. Europski revizorski sud. Dostupno na: https://www.eca.europa.eu/ECAPublications/SR-2024-08/SR-2024-08_HR.pdf [Pristupljeno: 30. 8. 2024.]
100. Faqih, K. M., & Jaradat, M. I. R. M. (2021). Integrating TTF and UTAUT2 theories to investigate the adoption of augmented reality technology in education: Perspective from a developing country. *Technology in Society*, 67, 101787.
101. Faraon, M., Rönkkö, K., Milrad, M., & Tsui, E. (2025). International perspectives on artificial intelligence in higher education: An explorative study of students' intention to use ChatGPT across the Nordic countries and the USA. *Education and Information Technologies*, 1–46.
102. Fishbein, M., & Ajzen, I. (1977). *Belief, attitude, intention, and behavior: An introduction to theory and research*.
103. Fitzmaurice, J. (2005). Incorporating consumers' motivations into the theory of reasoned action. *Psychology & Marketing*, 22 (11), 911–929.
104. Flasiński, M. (2016). *Introduction to artificial intelligence*. Switzerland: Springer International Publishing.
105. Floridi, L. (2020). AI and its new winter: From myths to realities. *Philosophy & Technology*, 33, 1–3.

106. Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of marketing research*, 18 (1), 39–50.
107. Foroughi, B., Senali, M. G., Iranmanesh, M., Khanfar, A., Ghobakhloo, M., Annamalai, N., & Naghmeh-Abbaspour, B. (2023). Determinants of intention to use ChatGPT for educational purposes: Findings from PLS-SEM and fsQCA. *International Journal of Human–Computer Interaction*, 1–20.
108. Future of Life (2024). *Pause Giant AI Experiments: An Open Letter*. Future of Life Institute. Dostupno na: <https://futureoflife.org/open-letter/pause-giant-ai-experiments/> [Pristupljeno: 3. 12. 2024.]
109. Gansser, O. A., & Reich, C. S. (2021). A new acceptance model for artificial intelligence with extensions to UTAUT2: An empirical study in three segments of application. *Technology in Society*, 65, 101535.
110. García de Blanes Sebastián M, Sarmiento Guede JR and Antonovica A (2022) Application and extension of the UTAUT2 model for determining behavioral intention factors in use of the artificial intelligence virtual assistants. *Front. Psychol.* 13:993935. doi: 10.3389/fpsyg.2022,993935.
111. Garner (2023). *Gartner Experts Answers the Top Generative AI Questions for Your Enterprise*. Gartner. Dostupno na: <https://www.gartner.com/en/topics/generative-ai> [Pristupljeno: 10. 9. 2024.]
112. Gefen, D. (2000). E-commerce: the role of familiarity and trust. *Omega*, 28 (6), 725–737.
113. Geisser, S. (1974). A Predictive Approach to the Random Effects Model, *Biometrika*, 61 (1): 101–107.
114. Goforth, C. (2015). *Using and Interpreting Cronbach's Alpha*. University of Virginia – StatLab Articles. Dostupno na: <https://library.virginia.edu/data/articles/using-and-interpreting-cronbachs-alpha> [Pristupljeno: 16. 4. 2025.]
115. Goldman Sachs (2023). *Generative AI could raise global GDP by 7%*. Goldman Sachs. Dostupno na: <https://www.goldmansachs.com/insights/articles/generative-ai-could-raise-global-gdp-by-7-percent.html> [Pristupljeno 1. 12. 2024.]

116. Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A. & Bengio, Y. (2014). Generative adversarial nets. *Advances in neural information processing systems*, 27.
117. Gøthesen, S., Haddara, M., & Kumar, K. N. (2023). Empowering homes with intelligence: An investigation of smart home technology adoption and usage. *Internet of Things*, 24, 100944.
118. Goyal, M., & Mahmoud, Q. H. (2024). A Systematic Review of Synthetic Data Generation Techniques Using Generative AI. *Electronics*, 13 (17), 3509.
119. Goyal, N., Kaur, H., & Mago, M. (2023) ChatGPT acceptance drivers: a study of university students in punjab. *Measurement*, 18 (4), 100.
120. Grassini, S., Aasen, M. L., & Møgelvang, A. (2024). Understanding University Students' Acceptance of ChatGPT: Insights from the UTAUT2 Model. *Applied Artificial Intelligence*, 38 (1), 2371168.
121. Greenland, S., Senn, S. J., Rothman, K. J., Carlin, J. B., Poole, C., Goodman, S. N., & Altman, D. G. (2016). Statistical tests, P values, confidence intervals, and power: a guide to misinterpretations. *European journal of epidemiology*, 31 (4), 337–350.
122. Gujarati, D. N. (2003). *Basic econometrics (4th ed.)*. McGraw-Hill.
123. Gursoy, D., Chi, O. H., Lu, L., & Nunkoo, R. (2019). Consumers acceptance of artificially intelligent (AI) device use in service delivery. *International Journal of Information Management*, 49, 157–169.
124. Habibi, A., Muhaimin, M., Danibao, B. K., Wibowo, Y. G., Wahyuni, S., & Octavia, A. (2023). ChatGPT in higher education learning: Acceptance and use. *Computers and Education: Artificial Intelligence*, 5, 100190.
125. Hagendorff, T. (2024). Mapping the ethics of generative ai: A comprehensive scoping review. *Minds and Machines*, 34 (4), 39.
126. Hair Jr, J. F., Howard, M. C., & Nitzl, C. (2020). Assessing measurement model quality in PLS-SEM using confirmatory composite analysis. *Journal of business research*, 109, 101–110.
127. Hair, J. F., Hult, G. T. M., Ringle, C. M. and Sarstedt, M. (2017b). *A Primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)*. 2. izd. Los Angeles: Sage Publications.

128. Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2022). *A primer on partial least squares structural equation modeling (PLS-SEM)* (3rd ed.). Sage Publications.
129. Hair, J. F., Hult, G. T. M., Ringle, C. M., Sarstedt, M., Danks, N. P., & Ray, S. (2021). *Partial least squares structural equation modeling (PLS-SEM) using R: A workbook* (p. 197). Springer Nature.
130. Hair, J. F., Matthews, L. M., Matthews, R. L. i Sarstedt, M. (2017a). PLS-SEM or CB-SEM: updated guidelines on which method to use. *International Journal of Multivariate Dana Analysis, Vol. 1* (2), str. 107–123.
131. Hair, J. F., Ringle, C. M. i Sarstedt, M. (2011). PLS-SEM: Indeed a Silver Bullet. *Journal of Marketing Theory and Practice, Vol. 19* (2), str. 139–151
132. Hair, J. F., Sarstedt, M., Pieper, T. M. i Ringle, C. M. (2012). The use of partial least squares structural equation modeling in strategic management research: a review of past practices and recommendations for future applications. *Long range planning, Vol. 45* (5-6), str. 320–340.
133. Hall, D. T., & Mansfield, R. (1975). Relationships of age and seniority with career variables of engineers and scientists. *Journal of Applied Psychology*, 60 (2), 201.
134. Hanafiah, M. H. (2020). Formative Vs. Reflective Measurement Model: Guidelines for Structural Equation Modeling Research. *International Journal of Analysis and Applications, Vol. 18* (5), str. 876–889.
135. Hanafizadeh, P., Behboudi, M., Koshksaray, A. A., & Tabar, M. J. S. (2014). Mobile-banking adoption by Iranian bank clients. *Telematics and informatics*, 31 (1), 62–78.
136. Hasher, L., and Zacks, R. T. (1979). Automatic and Effortful Processes in Memory, "*Journal of Experimental Psychology: General* (108:3), pp. 356–388
137. Hassan, M. S., Islam, M. A., Yusof, M. F. B., Nasir, H., & Huda, N. (2023). Investigating the determinants of Islamic mobile FinTech service acceptance: A modified UTAUT2 approach. *Risks*, 11 (2), 40.
138. Henseler, J., Hubona, G., & Ray, P. A. (2016). Using PLS path modeling in new technology research: Updated guidelines. *Industrial Management & Data Systems*, 116 (1), 2–20. <https://doi.org/10.1108/IMDS-09-2015-0382>.
139. Henseler, J., Ringle, C., & Sarstedt, M. (2015). A New Criterion for Assessing Discriminant Validity in Variance-based Structural Equation Modeling. *Journal of the*

Academy of Marketing Science, 43 (1), 115–135 <https://doi.org/10.1007/s11747-014-0403-8>

140. Hering, A. & Rojas, A. (2024). *AI at Work: Why GenAI is More Likely to Support Workers Than Replace Them*. Indeed. Dostupno na: <https://www.hiringlab.org/2024/09/25/artificial-intelligence-skills-at-work/> [Pristupljeno: 2. 12. 2024.]
141. Hew, J.-J., Lee, V.-H., Ooi, K.-B., & Wei, J. (2015). What catalyses mobile apps usage intention: an empirical analysis. *Industrial Management & Data Systems*, 115 (7), 1269–1291.
142. Hirschman, E. C. (1980). Innovativeness, novelty seeking, and consumer creativity. *Journal of consumer research*, 7 (3), 283–295.
143. Ho, J., Jain, A., & Abbeel, P. (2020). Denoising diffusion probabilistic models. *Advances in neural information processing systems*, 33, 6840–6851.
144. Holbrook, M. B., & Hirschman, E. C. (1982). The experiential aspects of consumption: Consumer fantasies, feelings, and fun. *Journal of consumer research*, 9 (2), 132–140.
145. Holbrook, M. B., & Hirschman, E. C. (1982). The experiential aspects of consumption: Consumer fantasies, feelings, and fun. *Journal of consumer research*, 9 (2), 132–140.
146. Hopp, T., & Gangadharbatla, H. (2016). Novelty effects in augmented reality advertising environments: The influence of exposure time and self-efficacy. *Journal of Current Issues & Research in Advertising*, 37 (2), 113–130.
147. Howarth, J. (2025). *Number of Parameters in GPT-4 (Latest Data)*. Exploding Topics. Dostupno na: <https://explodingtopics.com/blog/gpt-parameters> [Pristupljeno: 15. 5. 2025.].
148. Hu, K. (2023). *ChatGPT sets record for fastest-growing user base – analyst note*. Reuters. Dostupno na: <https://www.reuters.com/technology/chatgpt-sets-record-fastest-growing-user-base-analyst-note-2023-02-01/> [Pristupljeno: 3. 12. 2024.].
149. Hu, L. & Bentler, P.M. (1998). Fit indices in covariance structure modeling: sensitivity to underparameterized model misspecification, *Psychological Methods*, Vol. 3 No. 4, pp. 424–453.

150. Hu, L. & Bentler, P.M. (1999). Cutoff criteria for fit indexes in covariance structure analysis: conventional criteria versus new alternatives, *Structural Equation Modeling*, Vol. 6, No. 1, pp. 1–55.
151. Huang, M. H., & Rust, R. T. (2018). Artificial intelligence in service. *Journal of service research*, 21 (2), 155–172.
152. Huang, W., Ong, W. C., Wong, M. K. F., Ng, E. Y. K., Koh, T., Chandramouli, C., ... & Tromp, J. (2024). Applying the UTAUT2 framework to patients' attitudes toward healthcare task shifting with artificial intelligence. *BMC Health Services Research*, 24 (1), 455.
153. IBM (2020). *Artificial Intelligence (AI)*. IBM Cloud Learn Hub. Dostupno na: <https://www.ibm.com/uk-en/cloud/learn/what-is-artificial-intelligence> [Pristupljeno: 2. 12. 2024.].
154. IBM (2024). *What are large language models (LLM)?*. IBM. Dostupno na: <https://www.ibm.com/topics/large-language-models> [Pristupljeno: 2. 12. 2024.].
155. IMF (2024). *Indeks spremnosti na umjetnu inteligenciju*. Međunarodni monetarni fond – IMF. Dostupno na: <https://www.imf.org/external/datamapper/datasets/AIPI> [Pristupljeno: 27. 10. 2024.].
156. Jamaluddin, J., Abd Gaffar, N., & Din, N. S. S. (2023). Hallucination: A key challenge to Artificial Intelligence-Generated writing. *Malaysian Family Physician: the Official Journal of the Academy of Family Physicians of Malaysia*, 18, 68.
157. Jaspersen, J., Carter, P. E., & Zmud, R. W. (2005). A Comprehensive Conceptualization of the Post-Adoptive Behaviors Associated with IT-Enabled Work Systems. *Management Information Systems Quarterly*, 29 (3), 8.
158. Jiang, B. (2025). *DeepSeek's growth in China surges past ByteDance's Doubao in AI app race*. South China Morning Post. Dostupno na: <https://www.scmp.com/tech/big-tech/article/3297586/deepseeks-growth-china-surges-past-bytedances-doubao-ai-app-race> [Pristupljeno: 7. 2. 2025.].
159. Johnson, T. (2023). *Elon Musk Says “Something Good Will Come Of This” After Senate’s AI Forum, Chuck Schumer Signals AI Legislation Coming “In The General Category Of Months” — Update*. Deadline. Dostupno na: <https://deadline.com/2023/09/senate-ai-insight-forum-wga-elon-musk-mark-zuckerberg-1235545470/> [Pristupljeno: 3. 12. 2024.].

160. Jöreskog, K. G. (1971). Simultaneous factor analysis in several populations. *Psychometrika*, 36 (4), 409–426.
161. Jovanovic, M., & Campbell, M. (2022). Generative artificial intelligence: Trends and prospects. *Computer*, 55 (10), 107–112.
162. Kalota, F. (2024). A Primer on Generative Artificial Intelligence. *Education Sciences*, 14 (2), 172.
163. Kamal, M., & Subriadi, A. P. (2021, September). UTAUT model of mobile application: literature review. In *2021 International Conference on Electrical and Information Technology (IEIT)* (pp. 120-125). IEEE.
164. Kang, S., Yongjoo, C. & Kim, B. (2025). Impact of generative AI service adoption intent on user attitudes: Focusing on the Unified Theory of Acceptance and Use of Technology. *International Journal of Innovative Research and Scientific Studies*, 8 (1), 2021–2033
165. Kašparová, P. (2023). Intention to use business intelligence tools in decision making processes: applying a UTAUT 2 model. *Central European Journal of Operations Research*, 31 (3), 991–1008.
166. Kaya, F., Aydin, F., Schepman, A., Rodway, P., Yetişensoy, O., & Demir Kaya, M. (2024). The roles of personality traits, AI anxiety, and demographic factors in attitudes toward artificial intelligence. *International Journal of Human–Computer Interaction*, 40 (2), 497–514.
167. Khan, F. M., Singh, N., Gupta, Y., Kaur, J., Banik, S., & Gupta, S. (2022). A meta-analysis of mobile learning adoption in higher education based on unified theory of acceptance and use of technology 3 (UTAUT3). *Vision*, 09722629221101159.
168. Kilani, A. A. H. Z., Kakeesh, D. F., Al-Weshah, G. A., & Al-Debei, M. M. (2023). Consumer post-adoption of e-wallet: An extended UTAUT2 perspective with trust. *Journal of Open Innovation: Technology, Market, and Complexity*, 9(3), 100113.
169. Kim, S. S., Malhotra, N. K., & Narasimhan, S. (2005). Research note—two competing perspectives on automatic use: A theoretical and empirical comparison. *Information systems research*, 16 (4), 418–432.
170. Kingma, D. P. & Welling, M. (2013). Auto-encoding variational bayes. arXiv preprint arXiv:1312.6114.

171. Kirova, V. D., Ku, C. S., Laracy, J. R., & Marlowe, T. J. (2023). The ethics of artificial intelligence in the era of generative AI. *Journal of Systemics, Cybernetics and Informatics*, 21 (4), 42–50.
172. KMPG (2025). UK attitudes to AI. KMPG. Dostupno na: <https://kpmg.com/uk/en/insights/ai/uk-attitudes-to-ai.html> [Pristupljeno: 20. 5. 2025.].
173. Kock, N., & Lynn, G. S. (2012). Lateral collinearity and misleading results in variance-based SEM: An illustration and recommendations. *Journal of the Association for Information Systems*, 13 (7), 546–580. <https://doi.org/10.17705/1jais.00302>
174. Kononov, N. (2023). *ChatGPT as a digital assistant for startup entrepreneurs: Challenges, Opportunities and Perception*.
175. Kopal, R., Korkut, D., & Žnidar, K. (2025). Deep Insights into AI perception in Croatia. *Interdisciplinary Description of Complex Systems: INDECS*, 23 (1), 1–28.
176. Korkmaz, H., Fidanoglu, A., Ozcelik, S., & Okumus, A. (2022). User acceptance of autonomous public transport systems: Extended UTAUT2 model. *Journal of Public Transportation*, 24, 100013.
177. Krajcar, D. (n.d.). *Gari Kasparov u šahu pobijedio superračunalo Deep Blue*. Povijest.hr. Dostupno na: <https://povijest.hr/nadanasnjidan/garry-kasparov-u-sahu-pobijedio-superracunalo-deep-blue-1996/> [Pristupljeno: 2. 12. 2024.].
178. Kralj, L., Blažić, A., Valečić, H., Janeš, S., Blašković, V., Marinić, N., Slišurić, K., Dasović, D., Majdandžić, V. i Rakić, D. (2024). *Umjetna inteligencija u obrazovanju*. Agencija za elektroničke medije i UNICEF, Zagreb.
179. Kutner, M. H., Nachtsheim, C. J., Neter, J., & Li, W. (2005). *Applied linear statistical models*. McGraw-hill.
180. Kwa, T., West, B., Becker, J. Deng, A., Garcia, K., Hasin, M., Jawhar, S., Kinniment, M., Rush, N., Von Arx, S., Bloom, R., Broadley, T., Du, H., Goodrich, B., Jurković, N., Harold Miles, L., Nix, S., Lin, T., Parikh, N., Rein, D., Koba Sato, L.J., Wijk, H., Ziegler, D.M., Barnes, E. & Chan, L. (2025). *Measuring AI Ability to Complete Long Tasks*. METR. Dostupno na: <https://metr.org/blog/2025-03-19-measuring-ai-ability-to-complete-long-tasks/> [Pristupljeno: 20. 3. 2025.].
181. Lai, C. Y., Cheung, K. Y., Chan, C. S., & Law, K. K. (2024). Integrating the adapted UTAUT model with moral obligation, trust and perceived risk to predict

- ChatGPT adoption for assessment support: A survey with students. *Computers and Education: Artificial Intelligence*, 6, 100246.
182. Lai, P. C. (2017). The literature review of technology adoption models and theories for the novelty technology. *JISTEM-Journal of Information Systems and Technology Management*, 14 (1), 21–38.
 183. Lakhani, K. (2023). *AI won't replace humans – but humans with AI will replace humans without AI*. Harvard Business School. Dostupno na: <https://hbr.org/2023/08/ai-wont-replace-humans-but-humans-with-ai-will-replace-humans-without-ai> [Pristupljeno: 1. 12. 2024.].
 184. Lavidas, K., Voulgari, I., Papadakis, S., Athanassopoulos, S., Anastasiou, A., Filippidi, A., ... & Karacapilidis, N. (2024). Determinants of Humanities and Social Sciences Students' Intentions to Use Artificial Intelligence Applications for Academic Purposes. *Information*, 15 (6), 314.
 185. Law, M. (2024). *Sopra Steria: Gen AI to be a US\$100bn market by 2028*. Technology Magazine. Dostupno na: <https://technologymagazine.com/articles/sopra-steria-gen-ai-to-be-a-100bn-market-by-2028> [Pristupljeno: 15. 3. 2025.].
 186. Lee, C., Ward, C., Raue, M., D'Ambrosio, L., & Coughlin, J. F. (2017). Age differences in acceptance of self-driving cars: A survey of perceptions and attitudes. In *Human Aspects of IT for the Aged Population. Aging, Design and User Experience: Third International Conference, ITAP 2017, Held as Part of HCI International 2017, Vancouver, BC, Canada, July 9-14, 2017, Proceedings, Part I 3* (pp. 3-13). Springer International Publishing.
 187. Lee, H. J., Cho, H. J., Xu, W., & Fairhurst, A. (2010). The influence of consumer traits and demographics on intention to use retail self-service checkouts. *Marketing Intelligence & Planning*, 28 (1), 46–58.
 188. Lee, K. F., & Qiufan, C. (2021). *AI 2041: Ten visions for our future*. Crown Currency.
 189. Lee, K.F. (2021). *How AI will completely change the way we live in the next 20 years*. Time. Dostupno na: <https://time.com/6097625/kai-fu-lee-book-ai-2041/> [Pristupljeno: 20. 11. 2024.].
 190. Lee, S., Jones-Jang, S. M., Chung, M., Kim, N., & Choi, J. (2024). Who is using ChatGPT and why? Extending the Unified Theory of Acceptance and Use of

- Technology (UTAUT) model. *Information Research an international electronic journal*, 29 (1), 54–72.
191. Leschanowsky, A., Rech, S., Popp, B., & Bäckström, T. (2024). Evaluating privacy, security, and trust perceptions in conversational AI: A systematic review. *Computers in Human Behavior*, 108344.
 192. Levy, J. A. (1988). Intersections of gender and aging. *The Sociological Quarterly*, 29 (4), 479–486.
 193. Levy, S. (2024). 8 Google Employees Invented Modern AI. Here's the Inside Story. *Wired*. Dostupno na: <https://www.wired.com/story/eight-google-employees-invented-modern-ai-transformers-paper/> [Pristupljeno: 18. 6. 2025.].
 194. Li, L. (2010). A critical review of technology acceptance literature. *Referred Research Paper*, 4, 2010.
 195. Liker, J. K., & Sindi, A. A. (1997). User acceptance of expert systems: a test of the theory of reasoned action. *Journal of Engineering and Technology management*, 14 (2), 147–173.
 196. Limayem, M., & Hirt, S. G. (2003). Force of habit and information systems usage: Theory and initial validation. *Journal of the Association for information Systems*, 4 (1), 3.
 197. Limayem, M., Hirt, S. G., & Cheung, C. M. (2007). How habit limits the predictive power of intention: The case of information systems continuance. *MIS quarterly*, 705–737.
 198. Lin, L. & Parker, K. (2025). *Worker's exposure to AI*. Pew Research Center. Dostupno na: <https://www.pewresearch.org/social-trends/2025/02/25/workers-exposure-to-ai/> [Pristupljeno: 3. 3. 2025.].
 199. Liu, Y. & Wang, H. (2024). *Who on Earth Is Using Generative AI*. Digital Development Global Practice – World Bank Group. Dostupno na: <https://documents1.worldbank.org/curated/en/099720008192430535/pdf/IDU15f321eb5148701472d1a88813ab677be07b0.pdf> [Pristupljeno: 18. 12. 2024.].
 200. Liu, Y., Yang, Z., Yu, Z., Liu, Z., Liu, D., Lin, H., Li, M., Ma, S., & Shi, S. (2023). Generative artificial intelligence and its applications in materials science: Current situation and future perspectives. *Journal of Materiomics*, 9(4), 798-816.
 201. Lustig, C., Konkel, A., and Jacoby, L. L. 2004. "Which Route to Recovery?," *Psychological Science* (15:1 1), pp. 729–735.

202. Macedo, I. M. (2017). Predicting the acceptance and use of information and communication technology by older adults: An empirical examination of the revised UTAUT2. *Computers in human behavior*, 75, 935–948.
203. Madjarova, S. J., Williams III, R. J., Nwachukwu, B. U., Martin, R. K., Karlsson, J., Ollivier, M., & Pareek, A. (2022). Picking apart p values: common problems and points of confusion. *Knee Surgery, Sports Traumatology, Arthroscopy*, 30 (10), 3245–3248.
204. Maican, C. I., Sumedrea, S., Tecau, A., Nichifor, E., Chitu, I. B., Lixandriou, R., & Bratucu, G. (2023). Factors influencing the behavioural intention to use AI-Generated images in business: a UTAUT2 perspective with moderators. *Journal of Organizational and End User Computing (JOEUC)*, 35(1), 1-32.
205. Manyika, J. (2022). Getting AI right: Introductory notes on AI & society. *Daedalus*, 151 (2), 5–27.
206. Marikyan, D. & Papagiannidis, S. (2023). Unified Theory of Acceptance and Use of Technology: A review. In S. Papagiannidis (Ed), *TheoryHub Book*. Available at <https://open.ncl.ac.uk/> / ISBN: 9781739604400.
207. Markoff, J. (2005). *Behind Artificial Intelligence, a Squadron of Bright Real People*. New York Times. Dostupno na: <https://www.nytimes.com/2005/10/14/technology/behind-artificial-intelligence-a-squadron-of-bright-real-people.html> [Pristupljeno: 1. 12. 2024.].
208. Maruping, L. M., Bala, H., Venkatesh, V., & Brown, S. A. (2017). Going beyond intention: Integrating behavioral expectation into the unified theory of acceptance and use of technology. *Journal of the Association for Information Science and Technology*, 68 (3), 623–637.
209. Maslej, N., Fattorini, L., Perrault R., Parli V., Reuel A., Brynjolfsson E., Etchemendy J., Ligett K., Lyons T., Manyika J., Niebles J. C., Shoham Y., Wald R. & Clark J. (2024). *The AI Index 2024 Annual Report*. AI Index Steering Committee, Institute for Human-Centered Artificial Intelligence, Stanford University, Stanford, California, United States of America. Dostupno na: https://aiindex.stanford.edu/wp-content/uploads/2024/05/HAI_AI-Index-Report-2024.pdf [Pristupljeno: 27. 10. 2024.].
210. Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of management review*, 20 (3), 709–734.

211. Mays, K. K., Lei, Y., Giovanetti, R., & Katz, J. E. (2022). AI as a boss? A national US survey of predispositions governing comfort with expanded AI roles in society. *AI & Society*, 1–14.
212. McCarthy, J. (2004). *What is artificial intelligence*. Dostupno na: <https://www-formal.stanford.edu/jmc/whatisai/> [Pristupljeno: 30. 8. 2024.].
213. McCarthy, J., Minsky, M. L., Rochester, N., & Shannon, C. E. (1955). *A proposal for the dartmouth summer research project on artificial intelligence*. Stanford. Dostupno na: <http://jmc.stanford.edu/articles/dartmouth/dartmouth.pdf> [Pristupljeno: 30. 8. 2024.].
214. McCorduck, P., & Cfe, C. (2004). *Machines who think: A personal inquiry into the history and prospects of artificial intelligence*. AK Peters/CRC Press.
215. MDPI (2025). *Research and Publication Ethics – Artificial Intelligence*. MDPI. Dostupno na: <https://www.mdpi.com/ethics> [Pristupljeno: 13. 5. 2025.].
216. Meet, R. K., Kala, D., & Al-Adwan, A. S. (2022). Exploring factors affecting the adoption of MOOC in Generation Z using extended UTAUT2 model. *Education and Information Technologies*, 27 (7), 10261–10283.
217. Menon, D., & Shilpa, K. (2023). “Chatting with ChatGPT”: Analyzing the factors influencing users' intention to Use the Open AI's ChatGPT using the UTAUT model. *Heliyon*, 9 (11).
218. Merhi, M., Hone, K., & Tarhini, A. (2019). A cross-cultural study of the intention to use mobile banking between Lebanese and British consumers: Extending UTAUT2 with security, privacy and trust. *Technology in Society*, 59, 101151.
219. Microsoft (2024). *Global Online Safety Survey 2024*. Microsoft. Dostupno na: <https://news.microsoft.com/wp-content/uploads/prod/sites/40/2024/02/Microsoft-Global-Online-Safety-Survey-2024.pdf> [Pristupljeno: 18. 3. 2025.].
220. Midgley, D. F., & Dowling, G. R. (1978). Innovativeness: The concept and its measurement. *Journal of consumer research*, 4 (4), 229–242.
221. Miladinovic, J., & Hong, X. (2016). *A study on factors affecting the behavioral intention to use mobile shopping fashion apps in Sweden*.
222. Miller, J. B. (2012). *Toward a new psychology of women*. Beacon Press.
223. Minton, H. L., & Schneider, F. W. (1980). *Differential Psychology Waveland Press*. Prospect Heights.

224. Mohamad, M., Afthanorhan, A., Awang, Z. i Mohammad, M. (2019). Comparison Between CB-SEM and PLS-SEM: Testing and Confirming the Maqasid Syariah Quality of Life Measurement Model. *The Journal of Social Sciences Research*, Vol. 5 (3), str. 608–614
225. Mohd Rahim, N. I., A. Iahad, N., Yusof, A. F., & A. Al-Sharafi, M. (2022). AI-based chatbots adoption model for higher-education institutions: A hybrid PLS-SEM-neural network modelling approach. *Sustainability*, 14 (19), 12726.
226. Mondal, B. (2020). Artificial intelligence: state of the art. *Recent trends and advances in artificial intelligence and internet of things*, 389–425.
227. Moore, G. C., & Benbasat, I. (1991). Development of an instrument to measure the perceptions of adopting an information technology innovation. *Information systems research*, 2 (3), 192–222.
228. Moore, G. E. (1965). Cramming more components onto integrated circuits. *Electronics*, 38 (8).
229. Morgan, K., & Morgan, M. (2006). Gender-Biased Attitudes Toward Technology. In *Encyclopedia of Gender and Information Technology* (pp. 711-713). IGI Global.
230. Morris, M. G., Venkatesh, V., & Ackerman, P. L. (2005). Gender and age differences in employee decisions about new technology: An extension to the theory of planned behavior. *IEEE transactions on engineering management*, 52 (1), 69–84.
231. Mueller, M. (2024). *The Evolving Landscape of Large Language Model (LLM) Architectures*. Re-cinq. Dostupno na: <https://re-cinq.com/blog/llm-architectures> [Pristupljeno: 15. 5. 2025.].
232. Murugesan, S., & Cherukuri, A. K. (2023). The Rise of Generative Artificial Intelligence and Its Impact on Education: The Promises and Perils. *Computer*, 56(5), 116–121. <https://doi.org/10.1109/MC.2023.3253292>
233. National Research Council. (1999). *Funding a revolution: Government support for computing research*. National Academies Press.
234. Nikolopoulou, K., Gialamas, V., & Lavidas, K. (2020). Acceptance of mobile phone by university students for their studies: An investigation applying UTAUT2 model. *Education and Information Technologies*, 25, 4139–4155.

235. Nomura, T., Suzuki, T., Kanda, T., & Kato, K. (2006, September). Measurement of anxiety toward robots. In *ROMAN 2006-The 15th IEEE International Symposium on Robot and Human Interactive Communication* (pp. 372-377). IEEE.
236. Nordhoff, S., Louw, T., Innamaa, S., Lehtonen, E., Beuster, A., Torrao, G., ... & Merat, N. (2020). Using the UTAUT2 model to explain public acceptance of conditionally automated (L3) cars: A questionnaire study among 9, 118 car drivers from eight European countries. *Transportation research part F: Traffic psychology and behaviour*, 74, 280–297.
237. Nysveen, H., Pedersen, P. E., & Thorbjørnsen, H. (2005). Intentions to use mobile services: Antecedents and cross-service comparisons. *Journal of the academy of marketing science*, 33 (3), 330–346.
238. Oliveira, T., Thomas, M., Baptista, G., & Campos, F. (2016). Mobile payment: Understanding the determinants of customer adoption and intention to recommend the technology. *Computers in human behavior*, 61, 404–414.
239. Oliver Wyman Forum (2024). *How Generative AI is Transforming Business and Society: The Good, The Bad, and Everything in Between*. Oliver Wyman Forum. Dostupno na: <https://www.oliverwymanforum.com/global-consumer-sentiment/how-will-ai-affect-global-economics.html> [Pristupljeno: 3. 3. 2025.].
240. OpenAI (2024). ChatGPT-4 [Large Language Model]. Dostupno na: <https://chatgpt.com/> [Pristupljeno: 7. 12. 2024.].
241. Oxford Insights (2023). *Government AI Readiness Index 2023*. Oxford Insights. Dostupno na: <https://oxfordinsights.com/ai-readiness/ai-readiness-index/> [Pristupljeno: 27. 10. 2024.].
242. Pal, D., Roy, P., Arpnikanondt, C., & Thapliyal, H. (2022). The effect of trust and its antecedents towards determining users' behavioral intention with voice-based consumer electronic devices. *Heliyon*, 8(4).
243. Paul, J., Ueno, A., & Dennis, C. (2023). ChatGPT and consumers: Benefits, pitfalls and future research agenda. *International Journal of Consumer Studies*, 47 (4), 1213–1225.
244. Pichai, S. (2023). *An important next step on our AI journey*. Google. Dostupno na: <https://blog.google/technology/ai/bard-google-ai-search-updates/> [Pristupljeno: 3. 12. 2024.].

245. Pierce, D. (2023). *Google launches Gemini, the AI model it hopes will take down GPT-4*, The Verge. Dostupno na: <https://www.theverge.com/2023/12/6/23990466/google-gemini-llm-ai-model> [Pristupljeno: 3. 12. 2024.].
246. Plouffe, C. R., Hulland, J. S., & Vandenbosch, M. (2001). Richness versus parsimony in modeling technology adoption decisions—understanding merchant adoption of a smart card-based payment system. *Information systems research*, 12 (2), 208–222.
247. Plude, D. J. (1985). Attention and performance: Identifying and localizing age deficits. *Aging and Human Performance*, 47–99.
248. Porsdam Mann, S., Vazirani, A. A., Aboy, M., Earp, B. D., Minssen, T., Cohen, I. G., & Savulescu, J. (2024). Guidelines for ethical use and acknowledgement of large language models in academic writing. *Nature Machine Intelligence*, 1–3.
249. Porter, L. W. (1963). Job attitudes in management: II. Perceived importance of needs as a function of job level. *Journal of Applied Psychology*, 47 (2), 141.
250. Prizma (2024). *Percepcija umjetne inteligencije u RH 2024.; Pregled rezultata ispitivanja javnog mijenja*. Effectus. Dostupno na: <https://effectus.com.hr/wp-content/uploads/2024/10/2024-Izvjestaj-AI-sazeto.pdf> [Pristupljeno: 19. 2. 2025.].
251. Prolific (2025). Prolific. Dostupno na: <https://www.prolific.com/> [Pristupljeno: 19. 3. 2025.].
252. Rachini, M. (2022). *ChatGPT a 'landmark event' for AI, but what does it mean for the future of human labour and disinformation?*. CBC. Dostupno na: <https://www.cbc.ca/radio/thecurrent/chatgpt-human-labour-and-fake-news-1.6686210> [Pristupljeno: 3. 12. 2024.].
253. Rahman, M. M., Terano, H. J., Rahman, M. N., Salamzadeh, A., & Rahaman, M. S. (2023). ChatGPT and academic research: A review and recommendations based on practical examples. Rahman, M., Terano, HJR, Rahman, N., Salamzadeh, A., Rahaman, S.(2023). ChatGPT and Academic Research: A Review and Recommendations Based on Practical Examples. *Journal of Education, Management and Development Studies*, 3 (1), 1–12.
254. Reed, R. (2024). *ChatNYT*. Harvard Law Today. Dostupno na: <https://hls.harvard.edu/today/does-chatgpt-violate-new-york-times-copyrights/> [3. 12. 2024.].

255. Rhodes, S. R. (1983). Age-related differences in work attitudes and behavior: A review and conceptual analysis. *Psychological bulletin*, 93 (2), 328.
256. Ringle, Christian M., Wende, Sven, & Becker, Jan-Michael. (2024). *SmartPLS 4. Bönningstedt*: SmartPLS. Dostupno na: <https://www.smartpls.com> [Pristupljeno: 14. 4. 2025.].
257. Rios-Campos, C., Viteri, J. D. C. L., Batalla, E. A. P., Castro, J. F. C., Núñez, J. B., Calderón, E. V., ... & Tello, M. Y. P. (2023). Generative artificial intelligence. *South Florida Journal of Development*, 4 (6), 2305–2320.
258. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., & Ommer, B. (2022). High-resolution image synthesis with latent diffusion models. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition* (pp. 10684-10695).
259. Rumsey, D. J. (2016). *Statistics for dummies*. John Wiley & Sons.
260. Russell, S. J., & Norvig, P. (2021). *Artificial intelligence: a modern approach – Fourth Edition*. Pearson.
261. Ryan, R. M., & Deci, E. L. (2000). Self-determination theory and the facilitation of intrinsic motivation, social development, and well-being. *American Psychologist*, 55 (1), 68–78.
262. SAGE (2025). *Guidance on the use of generative artificial intelligence for authors submitting to the journal*, Language testing. SAGE. Dostupno na: <https://journals.sagepub.com/pb-assets/cmscontent/LTJ/LT%20Author%20guidelines%20on%20genAI-1723460922.pdf> [Pristupljeno: 13. 5. 2025.].
263. Salesforce (2024). *Top Generative AI Statistics for 2025*. Salesforce. Dostupno na: <https://www.salesforce.com/news/stories/generative-ai-statistics/#h-General-population-data-reveals-new-insights-> [Pristupljeno: 3. 3. 2025.].
264. Salifu, I., Arthur, F., Arkorful, V., Abam Nortey, S., & Solomon Osei-Yaw, R. (2024). Economics students' behavioural intention and usage of ChatGPT in higher education: A hybrid structural equation modelling-artificial neural network approach. *Cogent Social Sciences*, 10 (1), 2300177.
265. Sallam, M., Salim, N., Barakat, M., Al-Mahzoum, K., Al-Tammemi, A. B., Malaeb, D., ... & Hallit, S. (2023). Validation of a technology acceptance model-based

scale TAME-ChatGPT on health students attitudes and usage of ChatGPT in Jordan. *JMIR Preprints*.

266. Schepman, A., & Rodway, P. (2023). The General Attitudes towards Artificial Intelligence Scale (GAAIS): Confirmatory validation and associations with personality, corporate distrust, and general trust. *International Journal of Human–Computer Interaction*, 39 (13), 2724–2741.
267. Schmitz, A., Díaz-Martín, A. M., & Guillén, M. J. Y. (2022). Modifying UTAUT2 for a cross-country comparison of telemedicine adoption. *Computers in Human Behavior*, 130, 107183.
268. Senali, M. G., Iranmanesh, M., Ismail, F. N., Rahim, N. F. A., Khoshkam, M., & Mirzaei, M. (2023). Determinants of intention to use e-Wallet: Personal innovativeness and propensity to trust as moderators. *International Journal of Human–Computer Interaction*, 39 (12), 2361–2373.
269. Sengar, S. S., Hasan, A. B., Kumar, S., & Carroll, F. (2024). Generative Artificial Intelligence: A Systematic Review and Applications. arXiv preprint arXiv:2405.11029.
270. Sergeeva, O. V., Zheltukhina, M. R., Shoustikova, T., Tukhvatullina, L. R., Dobrokhoto, D. A., & Kondrashev, S. V. (2025). Understanding higher education students' adoption of generative AI technologies: An empirical investigation using UTAUT2. *Contemporary Educational Technology*, 17 (2), ep571.
271. Seseri, R., Martiro, K., & Dolce, K. (2024). *The History of Artificial Intelligence*. Glasswing Ventures. Dostupno na: <https://glasswing.vc/blog/thinking-corner/the-history-of-artificial-intelligence/> [Pristupljeno: 3. 10. 2024.].
272. Setiyani, L., Natalia, I., & Liswadi, G. T. (2023). Analysis of behavioral intentions of e-commerce shopee users in indonesia using utaut2. *ADI Journal on Recent Innovation*, 4 (2), 160–171.
273. Sharkey, N. (2012). *Alan Turing: The experiment that shaped artificial intelligence*. BBC. Dostupno na: <https://www.bbc.com/news/technology-18475646> [Pristupljeno: 2. 9. 2024.].
274. Sharma, R., Yetton, P., & Crawford, J. (2009). Estimating the effect of common method variance: the method—method pair technique with an illustration from TAM research. *Mis Quarterly*, 473–490.

275. Shaw, N., & Sergueeva, K. (2019). The non-monetary benefits of mobile commerce: Extending UTAUT2 with perceived value. *International journal of information management*, 45, 44–55.
276. Sheppard, B. H., Hartwick, J., & Warshaw, P. R. (1988). The theory of reasoned action: A meta-analysis of past research with recommendations for modifications and future research. *Journal of consumer research*, 15 (3), 325–343.
277. Shi, Y., Siddik, A. B., Masukujjaman, M., Zheng, G., Hamayun, M., & Ibrahim, A. M. (2022). The antecedents of willingness to adopt and pay for the IoT in the agricultural industry: an application of the UTAUT 2 theory. *Sustainability*, 14 (11), 6640.
278. Sinaga, J. N., Panjaitan, E. S., & Nurjanah, S. (2024). Analysis of Factors Affecting the Use of ChatGPT at Mikroskil University: A Study Based on the Extended UTAUT2 Model. *Brilliance: Research of Artificial Intelligence*, 4 (1), 151–161.
279. Singla, A., Sukharevsky, A., Yee, L., Chui, M. & Hall, B. (2024). *The state of AI in early 2024: Gen AI adoption spikes and starts to generate value*. McKinsey. Dostupno na: <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai> [Pristupljeno: 4. 12. 2024.].
280. Sobaih, A. E. E., Elshaer, I. A., & Hasanein, A. M. (2024). Examining Students' Acceptance and Use of ChatGPT in Saudi Arabian Higher Education. *European Journal of Investigation in Health, Psychology and Education*, 14 (3), 709–721.
281. Sohn, K., & Kwon, O. (2020). Technology acceptance theories and factors influencing artificial Intelligence-based intelligent products. *Telematics and Informatics*, 47, 101324.
282. Soman, D., & Cheema, A. (2001). The effect of windfall gains on the sunk-cost effect. *Marketing Letters*, 12, 51–62.
283. Soman, D., & Gourville, J. T. (2001). Transaction decoupling: How price bundling affects the decision to consume. *Journal of marketing research*, 38 (1), 30–44.
284. Springer (2025). *Editorial policies – Artificial intelligence (AI)*. Springer. Dostupno na: <https://www.springer.com/gp/editorial-policies/artificial-intelligence--ai->

[/25428500?srsId=AfmBOoqbBQGsrW9gCaAE2YqMkQvzAiL_27L43Az2hnbdfqAV2EIMcYK0](https://doi.org/10.1007/978-1-4939-9966-6_25428500?srsId=AfmBOoqbBQGsrW9gCaAE2YqMkQvzAiL_27L43Az2hnbdfqAV2EIMcYK0) [Prisutpljeno: 13. 5. 2025.].

285. Springer Nature (2025). *AI Principles*. Springer Nature. Dostupno na: <https://group.springernature.com/gp/group/ai/ai-principles> [Prisutpljeno: 13. 5. 2025.].
286. Središnji državni ured za razvoj digitalnog društva (2023). *Strategija digitalne Hrvatske za razdoblje do 2032*. Republika Hrvatska. Dostupno na: https://rdd.gov.hr/UserDocsImages//SDURDD-dokumenti//Strategija_digitalne_Hrvatske_za_razdoblje_do_2032.pdf [Prisutpljeno: 27. 10. 2024.].
287. Statista (2024). *Generative AI*. Statista. Dostupno na: <https://www.statista.com/outlook/tmo/artificial-intelligence/worldwide> [Prisutpljeno: 18. 3. 2025.].
288. Sterling, B. (2020). *Web Semantics: Microsoft Project Turing introduces Turing Natural Language Generation (T-NLG)*. Wired. Dostupno na: <https://www.wired.com/beyond-the-beyond/2020/02/web-semantics-microsoft-project-turing-introduces-turing-natural-language-generation-t-nlg/> [Prisutpljeno: 3. 12. 2024.].
289. Stone, M. (1974). Cross-Validatory Choice and Assessment of Statistical Predictions, *Journal of the Royal Statistical Society*, 36 (2): pp 111-147.
290. Stryker, C. & Scapicchio, M. (2024). *What is generative AI?*, IBM. Dostupno na: <https://www.ibm.com/topics/generative-ai> [Prisutpljeno: 11. 9. 2024.].
291. Strzelecki, A. (2023). To use or not to use ChatGPT in higher education? A study of students' acceptance and use of technology. *Interactive learning environments*, 1–14.
292. Sugumar, M., & Chandra, S. (2021). Do I desire chatbots to be like humans? Exploring factors for adoption of chatbots for financial services. *Journal of international technology and information management*, 30 (3), 38–77.
293. Sun, W., Shin, H. Y., Wu, H., & Chang, X. (2023). Extending UTAUT2 with knowledge to test Chinese consumers' adoption of imported spirits flash delivery applications. *Heliyon*, 9 (5).
294. Svjetski ekonomski forum (2024). *Women are falling behind on generative AI in the workplace: Here's how to change that*. World Economic Forum. Dostupno na:

<https://www.weforum.org/stories/2024/04/women-generative-ai-workplace/>

[Pristupljeno: 3. 3. 2025.].

295. Tamilmani, K., Rana, N. P., & Dwivedi, Y. K. (2017). A systematic review of citations of UTAUT2 article and its usage trends. In *Digital Nations–Smart Cities, Innovation, and Sustainability: 16th IFIP WG 6,11 Conference on e-Business, e-Services, and e-Society, I3E 2017, Delhi, India, November 21–23, 2017, Proceedings 16* (pp. 38-49). Springer International Publishing.
296. Tamilmani, K., Rana, N. P., & Dwivedi, Y. K. (2017). A systematic review of citations of UTAUT2 article and its usage trends. In *Digital Nations–Smart Cities, Innovation, and Sustainability: 16th IFIP WG 6,11 Conference on e-Business, e-Services, and e-Society, I3E 2017, Delhi, India, November 21–23, 2017, Proceedings 16* (pp. 38-49). Springer International Publishing.
297. Tamilmani, K., Rana, N. P., & Dwivedi, Y. K. (2021). Consumer acceptance and use of information technology: A meta-analytic evaluation of UTAUT2. *Information Systems Frontiers*, 23, 987–1005.
298. Tamilmani, K., Rana, N. P., Wamba, S. F., & Dwivedi, R. (2021). The extended Unified Theory of Acceptance and Use of Technology (UTAUT2): A systematic literature review and theory evaluation. *International Journal of Information Management*, 57, 102269.
299. Tams, S., & Dulipovici, A. (2024). The creativity model of age and innovation with IT: why older users are less innovative and what to do about it. *European Journal of Information Systems*, 33 (3), 287–314.
300. Tantra, T., & Ariyanti, M. (2017, November). The Use of Modified Unified Theory of Acceptance and Use of Technology 2 (UTAUT2) to predict Student Behavioral Intention in the use of Integrated Academic Information System (iGracias) Mobile Application at Telkom University. In *3rd International Conference on Transformation in Communications 2017 (IcoTiC 2017)* (pp. 96-101). Atlantis Press.
301. Tarhini, A., AlHinai, M., Al-Busaidi, A. S., Govindaluri, S. M., & Al Shaqsi, J. (2024). What drives the adoption of mobile learning services among college students: An application of SEM-neural network modeling. *International Journal of Information Management Data Insights*, 4 (1), 100235.
302. Taylor & Francis (2025). *AI Policy*. Taylor & Francis. Dostupno na: <https://taylorandfrancis.com/our-policies/ai-policy/> [Pristupljeno: 13. 5. 2025.].

303. Taylor, A. (2024). *IDC's 2024 AI opportunity study: Top five AI trends to watch*. *Microsoft Blog*. Dostupno na: <https://blogs.microsoft.com/blog/2024/11/12/idcs-2024-ai-opportunity-study-top-five-ai-trends-to-watch/> [29. 11. 2024.].
304. Taylor, S., & Todd, P. (1995b). Assessing IT usage: The role of prior experience. *MIS quarterly*, 561–570.
305. Taylor, S., & Todd, P. A. (1995a). Understanding information technology usage: A test of competing models. *Information systems research*, 6 (2), 144–176.
306. The White House (2023). *Executive Order on Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence*. White House. Dostupno na: <https://www.whitehouse.gov/briefing-room/presidential-actions/2023/10/30/executive-order-on-the-safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence/> [Pristupljeno: 23. 10. 2024.].
307. Thompson, R. L., Higgins, C. A., & Howell, J. M. (1991). Personal computing: Toward a conceptual model of utilization. *MIS quarterly*, 125–143.
308. Thusi, P., & Maduku, D. K. (2020). South African millennials' acceptance and use of retail mobile banking apps: An integrated perspective. *Computers in human behavior*, 111, 106405.
309. Tseng, T. H., Lin, S., Wang, Y. S., & Liu, H. X. (2022). Investigating teachers' adoption of MOOCs: the perspective of UTAUT2. *Interactive Learning Environments*, 30 (4), 635–650.
310. Turing, A. (1995). "Computing machinery and intelligence.". *Mind*, 59.
311. Tyson, L. (2022). *The Decade of AI Development: The Most Noteworthy Moments of the 2010s*. Medium. Dostupno na: <https://medium.com/@lenatyson/the-decade-of-ai-development-the-most-noteworthy-moments-of-the-2010s-983d2f299d49> [Pristupljeno: 15. 5. 2025.].
312. UN (2024). *Lethal Autonomous Weapon Systems (LAWS)*. United Nations. Dostupno na: <https://disarmament.unoda.org/the-convention-on-certain-conventional-weapons/background-on-laws-in-the-ccw/> [Pristupljeno: 30. 11. 2024.].
313. University of Oxford (2024). *New ethical framework to help navigate use of AI in academic research*. University of Oxford. Dostupno na: <https://www.ox.ac.uk/news/2024-11-13-new-ethical-framework-help-navigate-use-ai-academic-research> [Pristupljeno: 13. 5. 2025.].

314. Van der Heijden, H. (2004). User acceptance of hedonic information systems. *MIS quarterly*, 695–704.
315. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, U., & Polosukhin, I. (2017). Attention is all you need, *Advances in neural information processing systems* 30 (pp. 5999–6009). Curran Associates Inc.
316. Venkatesh, V. (2000). Determinants of perceived ease of use: Integrating control, intrinsic motivation, and emotion into the technology acceptance model. *Information systems research*, 11 (4), 342–365.
317. Venkatesh, V. (2022). Adoption and use of AI tools: a research agenda grounded in UTAUT. *Annals of Operations Research*, 308(1), 641-652.
318. Venkatesh, V. (2024). *Službena stranica – o meni*. Vvenkatesh. Dostupno na: <https://www.vvenkatesh.com/about/short-bio> [Pristupljeno: 4. 1. 2024.].
319. Venkatesh, V., & Bala, H. (2008). Technology acceptance model 3 and a research agenda on interventions. *Decision sciences*, 39 (2), 273–315.
320. Venkatesh, V., & Davis, F. D. (2000). A theoretical extension of the technology acceptance model: Four longitudinal field studies. *Management science*, 46 (2), 186–204.
321. Venkatesh, V., Brown, S. A., Maruping, L. M., & Bala, H. (2008). Predicting different conceptualizations of system use: The competing roles of behavioral intention, facilitating conditions, and behavioral expectation. *MIS quarterly*, 483–502.
322. Venkatesh, V., Morris, M. G., & Ackerman, P. L. (2000). A longitudinal field investigation of gender differences in individual technology adoption decision-making processes. *Organizational behavior and human decision processes*, 83 (1), 33–60.
323. Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS quarterly*, 425–478.
324. Venkatesh, V., Thong, J. Y., & Xu, X. (2012). Consumer acceptance and use of information technology: extending the unified theory of acceptance and use of technology. *MIS quarterly*, 157–178.
325. Venkatesh, V., Thong, J. Y., & Xu, X. (2016). Unified theory of acceptance and use of technology: A synthesis and the road ahead. *Journal of the association for Information Systems*, 17 (5), 328–376.

326. Vimalkumar, M., Sharma, S. K., Singh, J. B., & Dwivedi, Y. K. (2021). 'Okay google, what about my privacy?': User's privacy perceptions and acceptance of voice based digital assistants. *Computers in Human Behavior*, 120, 106763.
327. Vinerean, S., Budac, C., Baltador, L. A., & Dabija, D. C. (2022). Assessing the effects of the COVID-19 pandemic on M-commerce adoption: an adapted UTAUT2 approach. *Electronics*, 11 (8), 1269.
328. Vlada UK (2024). *Grow Your Tech Business in the UK*. UK Government. Dostupno na: <https://www.great.gov.uk/campaign-site/grow-your-tech-business-in-the-uk/artificial-intelligence/> [Pristupljeno: 17. 5. 2025.]
329. Voorhees, C. M., Brady, M. K., Calantone, R., & Ramirez, E. (2015). Discriminant validity testing in marketing: an analysis, causes for concern, and proposed remedies. *Journal of the Academy of Marketing Science*, 44 (1), 119–134. doi:10.1007/s11747-015-0455-4
330. Vuković, M. (2022). *Strukturalno modeliranje utjecaja bihevioralnih faktora na odlučivanje i performanse investitora na financijskom tržištu* (Doctoral dissertation, University of Split. Faculty of economics Split).
331. Wang, G. (2024). *LLMs milestons*. Medium. Dostupno na: <https://medium.com/@gremwang/llms-milestones-573e66737577> [Pristupljeno: 15. 5. 2025.]
332. Wei, R., & Mahmood, A. (2020). Recent advances in variational autoencoders with representation learning for biomedical informatics: A survey. *IEEE Access*, 9, 4939–4956.
333. Weitzman, T. (2023). *GPT-4 Released: What It Means For The Future Of Your Business*. Forbes. Dostupno na: <https://www.forbes.com/councils/forbesbusinesscouncil/2023/03/28/gpt-4-released-what-it-means-for-the-future-of-your-business/> [Pristupljeno: 3. 12. 2024.]
334. Wells, J. D., Campbell, D. E., Valacich, J. S., & Featherman, M. (2010). The effect of perceived novelty on the adoption of information technology innovations: a risk/reward perspective. *Decision Sciences*, 41 (4), 813–843.
335. Wiemer, M. (2023). *ChatGPT was never just a language model, despite OpenAI's claims*. Medium. Dostupno na: <https://markwiemer.medium.com/chatgpt->

[was-never-just-a-language-model-despite-openais-claims-60778d81eb04](#)

[Pristupljeno: 2. 12. 2024.].

336. Wiley (2025). *Using AI tools in your writing*. Wiley. Dostupno na: <https://www.wiley.com/en-us/publish/book/ai-guidelines> [Pristupljeno: 13. 5. 2025.].
337. WIPO (2024). *World Intellectual Property Organization. Generative AI: The main concepts*. WIPO. Dostupno na: <https://www.wipo.int/web-publications/patent-landscape-report-generative-artificial-intelligence-genai/en/1-generative-ai-the-main-concepts.html> [Pristupljeno: 2. 12. 2024.].
338. Xian, X. (2021). Psychological factors in consumer acceptance of artificial intelligence in leisure economy: a structural equation model. *Journal of Internet Technology*, 22 (3), 697–705.
339. Xu, J., Li, Y., Shadiev, R., & Li, C. (2025). College students' use behavior of generative AI and its influencing factors under the unified theory of acceptance and use of technology model. *Education and Information Technologies*, 1–24.
340. Xu, S., Chen, P., & Zhang, G. (2024). Exploring Chinese University Educators' Acceptance and Intention to use AI tools: An application of the UTAUT2 model. *SAGE Open*, 14 (4), 21582440241290013.
341. Zao-Sanders, M. (2025). *How People Are Really Using Gen AI in 2025*. Harvard Business Review. Dostupno na: <https://hbr.org/2025/04/how-people-are-really-using-gen-ai-in-2025> [Pristupljeno: 20. 6. 2025.].
342. Zhang, S., McClean, S. I., Nugent, C. D., Donnelly, M. P., Galway, L., Scotney, B. W., & Cleland, I. (2013). A predictive model for assistive technology adoption for people with dementia. *IEEE journal of biomedical and health informatics*, 18 (1), 375–383.
343. Zhang, X., Guo, X., Lai, K. H., Guo, F., & Li, C. (2014). Understanding gender differences in m-health adoption: a modified theory of reasoned action model. *Telemedicine and e-Health*, 20 (1), 39–46.

Prilog 1 - Oznake čestica

Konstrukt	Čestica	Kod
Namjera ponašanja	<i>Namjeravam nastaviti koristiti alate generativne umjetne inteligencije u svakodnevnom životu</i>	BI1
	<i>Uvijek ću nastojati koristiti alate generativne umjetne inteligencije u svakodnevnom životu</i>	BI2
	<i>Planiram često koristiti alate generativne umjetne inteligencije</i>	BI3
Očekivani učinak	<i>Smatram da je korištenje alata generativne umjetne inteligencije korisno</i>	PE1
	<i>Korištenje alata generativne umjetne inteligencije povećava moje šanse da ostvarim stvari koje su mi bitne</i>	PE2
	<i>Korištenje alata generativne umjetne inteligencije pomaže mi da brže obavim aktivnosti</i>	PE3
	<i>Korištenje alata generativne umjetne inteligencije povećava moju produktivnost</i>	PE4
Očekivani naponi	<i>Lako mi je naučiti koristiti alate generativne umjetne inteligencije</i>	EE1
	<i>Moja interakcija s alatima generativne umjetne inteligencije je jasna i razumljiva</i>	EE2
	<i>Smatram da su alati generativne umjetne inteligencije jednostavni za korištenje</i>	EE3
	<i>Lako mi je postati vješt u korištenju generativne umjetne inteligencije</i>	EE4
Društveni utjecaj	<i>Osobe koje su mi važne misle da bih trebao koristiti alate generativne umjetne inteligencije</i>	SI1
	<i>Osobe koje utječu na moje ponašanje misle da bih trebao koristiti alate generativne umjetne inteligencije</i>	SI2
	<i>Osobe čije mišljenje cijenim preferiraju da koristim alate generativne umjetne inteligencije</i>	SI3
Olakšavajući uvjeti	<i>Imam potrebne resurse za korištenje alata generativne umjetne inteligencije</i>	FC1
	<i>Imam potrebno znanje za korištenje generativne umjetne inteligencije</i>	FC2
	<i>Alati generativne umjetne inteligencije kompatibilni su s drugim tehnologijama koje koristim</i>	FC3
	<i>Mogu dobiti pomoć od drugih kada imam poteškoća s korištenjem alata generativne umjetne inteligencije</i>	FC4
Hedonistička motivacija	<i>Korištenje generativne umjetne inteligencije je zabavno</i>	HM1
	<i>Korištenje generativne umjetne inteligencije je ugodno</i>	HM2
	<i>Korištenje generativne umjetne inteligencije je vrlo zanimljivo</i>	HM3
Cijena	<i>Alati generativne umjetne inteligencije pružaju dobru vrijednost za novac</i>	PV1
	<i>Alati generativne umjetne inteligencije imaju razumnu cijenu</i>	PV2

	<i>Po trenutnoj cijeni, alati generativne umjetne inteligencije pružaju dobru vrijednost</i>	PV3
Navika	<i>Korištenje alata generativne umjetne inteligencije postalo mi je navika</i>	HT1
	<i>Ovisan sam o korištenju alata generativne umjetne inteligencije</i>	HT2
	<i>Moram koristiti alate generativne umjetne inteligencije</i>	HT3
	<i>Korištenje alata generativne umjetne inteligencije postalo mi je prirodno</i>	HT4
Osobna inovativnost	<i>Kada čujem za novu tehnologiju potražim način da ju odmah isprobam</i>	PI1
	<i>Među mojim prijateljima, ja sam uglavnom prvi/a koji proba nove tehnologije</i>	PI2
	<i>Volim eksperimentirati s novim tehnologijama</i>	PI3
Povjerenje	<i>Koristit ću alate generativne umjetne inteligencije ako osjetim da je sadržaj pouzdan</i>	TR1
	<i>Koristit ću alate generativne umjetne inteligencije ako osjetim da mi pruža pouzdane informacije</i>	TR2
	<i>Koristit ću alate generativne umjetne inteligencije ako osjetim da ispunjava moja očekivanja</i>	TR3
	<i>Koristit ću alate generativne umjetne inteligencije ako osjetim da je sigurna</i>	TR4
Korištenje	<i>U prosjeku, koliko sati tjedno koristite alate generativne umjetne inteligencije?</i>	USE1
	<i>Koliko često koristite alate generativne umjetne inteligencije?</i>	USE2
	<i>Kako procjenjujete intenzitet Vaše trenutačne uporabe alata generativne umjetne inteligencije?</i>	USE3

11. POPIS TABLICA

Tablica 1. Najkorišteniji alati generativne umjetne inteligencije	41
Tablica 2. Korijen konstrukta – Očekivani učinak (PE).....	54
Tablica 3. Korijen konstrukta – Očekivani naponi (EE).....	58
Tablica 4. Korijen konstrukta – društveni utjecaj (SI)	60
Tablica 5. Korijeni konstrukta – olakšavajući uvjeti (FC)	62
Tablica 6. Korijen konstrukta – hedonističke motivacije (HM)	67
Tablica 7. Korijen konstrukta – vrijednosti cijene (PV).....	68
Tablica 8. Korijen konstrukta – navika (HT)	68
Tablica 9. Pregled literature	76
Tablica 10. Korišteni standardni konstrukti iz UTAUT2 modela	89
Tablica 11. Najčešće korišteni konstrukti kao proširenja modela za istraživanje prihvatanja i korištenja tehnologije	90
Tablica 12. Varijacije čestica za konstrukt stvarnog korištenja	92
Tablica 13. Pregled radova koji se bave generativnom umjetnom inteligencijom.....	107
Tablica 14. Usporedba Republike Hrvatske i Ujedinjenog Kraljevstva po dimenzijama Indeksa spremnosti na umjetnu inteligenciju.....	124
Tablica 15. Istraživački instrument na hrvatskom i engleskom jeziku.....	129
Tablica 16. Preporuke veličine uzorka za PLS-SEM pri snazi testa od 80 %.....	141
Tablica 17. Prikaz kvotnog uzorka na 400 ispitanika.....	142
Tablica 18. Demografske karakteristike cjelokupnog uzorka	143
Tablica 19. Statistički pokazatelji i validacija mjernih varijabli [Cjeloviti uzorak].....	146
Tablica 20. Unakrsna faktorska opterećenja konstrukta [Cjeloviti uzorak]	149
Tablica 21. Pokazatelji unutarnje konzistentnosti i konvergentne valjanosti [Cjeloviti uzorak]	151
Tablica 22. Diskriminantna valjanost prema Fornell-Larckerovom kriteriju [Cjeloviti uzorak]	152
Tablica 23. Procjena diskriminantne valjanosti korištenjem HTMT kriterija [Cjeloviti uzorak]	152
Tablica 24. Rezultati testiranja hipoteza istraživanja [Cjeloviti uzorak].....	156
Tablica 25. Testiranje odnosa u modelu bez moderatora [Cjeloviti uzorak].....	160
Tablica 26. Demografske karakteristike uzorka na tržištu Ujedinjenog Kraljevstva	161
Tablica 27. Statistički pokazatelji i validacija mjernih varijabli [UK uzorak]	163
Tablica 28. Unakrsna faktorska opterećenja konstrukta [UK uzorak]	166
Tablica 29. Pokazatelji unutarnje konzistentnosti i konvergentne valjanosti [UK uzorak].....	168
Tablica 30. Diskriminantna valjanost prema Fornell-Larckerovom kriteriju [UK uzorak]	169
Tablica 31. Procjena diskriminantne valjanosti korištenjem HTMT kriterija [UK uzorak].....	169
Tablica 32. Rezultati testiranja odnosa u konceptualnom modelu [UK uzorak].....	171
Tablica 33. Testiranje odnosa u modelu bez moderatora [UK uzorak]	174
Tablica 34. Demografske karakteristike uzorka na tržištu Republike Hrvatske	176
Tablica 35. Statistički pokazatelji i validacija mjernih varijabli [RH uzorak]	178
Tablica 36. Unakrsna faktorska opterećenja konstrukta [RH uzorak].....	180
Tablica 37. Pokazatelji unutarnje konzistentnosti i konvergentne valjanosti [RH uzorak]	182
Tablica 38. Diskriminantna valjanost prema Fornell-Larckerovom kriteriju [RH uzorak].....	183
Tablica 39. Procjena diskriminantne valjanosti korištenjem HTMT kriterija [RH uzorak]	183
Tablica 40. Rezultati testiranja odnosa u konceptualnom modelu [RH uzorak]	185
Tablica 41. Testiranje odnosa u modelu bez moderatora (RH uzorak)	188

12. POPIS SLIKA

Slika 1. Prikaz bitnih događaja u razvoju umjetne inteligencije	23
Slika 2. Primjene i polja umjetne inteligencije.....	32
Slika 3. Pojednostavljeni prikaz funkcioniranja velikih jezičnih modela.....	37
Slika 4. Prikaz modela – Teorija razumnog djelovanja.....	48
Slika 5. Model prihvaćanja tehnologije (TAM)	48
Slika 6. Prikaz modela – Model prihvaćanje tehnologije 2 (TAM2)	49
Slika 7. Prikaz razvoja modela TAM 3	50
Slika 8. Prikaz modela – Teorija planiranog ponašanja TPB	51
Slika 9. Prikaz modela – kombinacija Modela prihvaćanja tehnologije i Teorije planiranog ponašanja (C-TAM-TPB).....	52
Slika 10. Prikaz modela – Model korištenja osobnih računala.....	53
Slika 11. Prikaz UTAUT modela	54
Slika 12. Prikaz modela UTAUT2 Izvor: prilagođeno prema Venkateshu i sur. (2012)	69
Slika 13. Pregled literature – prikaz utjecaja među varijablama koji su potvrđeni	110
Slika 14. Pregled literature – prikaz utjecaja među varijablama koji nisu potvrđeni	115
Slika 15. Karta svijeta po IMF-ovom Indeksu spremnosti na umjetnu inteligenciju	123
Slika 16. Prikaz konceptualnog modela – prošireni UTAUT2 model Izvor: izrada autora.....	135
Slika 17. Razlike radno aktivnih korisnika generativne umjetne inteligencije po dobi i spolu.....	142
Slika 18. Rezultati empirijskog istraživanja na konceptualnom modelu [Cjeloviti uzorak]	159
Slika 19. Rezultati empirijskog istraživanja na konceptualnom modelu [UK uzorak]	173
Slika 20. Rezultati empirijskog istraživanja na konceptualnom modelu [RH uzorak].....	187
Slika 21. Empirijski rezultati na UTAUT modelu.....	205
Slika 22. Empirijski rezultati na UTAUT2 modelu.....	206
Slika 23. Predloženi skraćeni model	211

13. POPIS KRATICA

UK – Ujedinjeno Kraljevstvo / Ujedinjeno Kraljevstvo Velike Britanije i Sjeverne Irske

AI – Umjetna inteligencija (engl. *Artificial intelligence*)

GEN AI – Generativna umjetna inteligencija (engl. *Generative artificial intelligence*)

RH – Republika Hrvatska

PE – Očekivani učinci (engl. – *Performance Expectancy*)

EE – Očekivani naponi (engl. – *Effort Expectancy*)

SI – Društveni utjecaj (engl. *Social influence*)

FC – Olakšavajući uvjeti (engl. *Facilitating conditions*)

HM – Hedonistička motivacija (engl. *Hedonic motivation*)

PV – Vrijednost cijene / cijena (engl. *Price value*)

HT – Navika (engl. *Habit*)

PI – Osobna inovativnost (engl. *Personal innovativeness*)

TR – Povjerenje (engl. *Trust*)

BI – Namjera ponašanja (engl. *Behavioral intention*)

USE – Korištenje (engl. *Use / Usage / Use behavior*)

UTAUT – Ujedinjena teorija prihvaćanja i korištenja tehnologije (engl. *Unified Theory of Acceptance and Use of Technology*)

UTAUT2 - Ujedinjena teorija prihvaćanja i korištenja tehnologije 2 tehnologije (engl. *Unified Theory of Acceptance and Use of Technology 2*)

TAM – Model prihvaćanja tehnologije (engl. *Technology acceptance model*)

TAM – Model prihvaćanja tehnologije (engl. *Technology acceptance model*)

IMF – Međunarodni monetarni fond (engl. *International Monetary Fund*)

LLM – Veliki jezični modeli (engl. *Large-language models*)

PLS-SEM – Modeliranje strukturnih jednadžbi djelomičnih najmanjih kvadrata (engl. *Partial least squares path modeling*)

TRA – Teorija razumnog djelovanja (engl. *Theory of reasoned action*)

MM – Motivacijski model (engl. *Motivational model*)

TPB – Teorija planiranog ponašanja (engl. *Theory of Planned Behavior*)

C-TAM-TPB – Kombinirani TAM i TPB (engl. *Combined TAM and TPB*)

MPCU – Model korištenja osobnih računala (engl. *Model of PC Utilization*)

IDT – Teorija difuzije inovacija (engl. *Innovation Diffusion Theory*)

SCT - Socijalna kognitivna teorija (engl. *Social Cognitive Theory*)

14. POPIS OBJAVLJENIH RADOVA PRISTUPNIKA

Biloš, A., **Budimir, B.**, & Jaška, S. (2021). Pozicija i značaj influencera u hrvatskoj. *CroDiM: International Journal of Marketing Science*, 4 (1), 57–68.

Biloš, A., & **Budimir, B.** (2024). Understanding the adoption dynamics of ChatGPT among generation Z: Insights from a modified UTAUT2 model. *Journal of theoretical and applied electronic commerce research*, 19 (2), 863–879.

Biloš, A., **Budimir, B.**, & Hrustek, A. (2022). The role of user-generated content in tourists' travel planning behavior: evidence from Croatia. *Revista Turismo & Desenvolvimento (RT&D)/Journal of Tourism & Development*, (39).

Ružić, D., Biloš, A., & **Budimir, B.** (2017). Exploring the influencing factors on the perception of web-shop customers in Croatia: a preliminary study. *Customer relationship management The impact of digital technology*, 23–34.

Biloš, A., **Budimir, B.**, & Kraljević, B. (2023). Attitudes and preferences toward the adoption of voice-controlled intelligent personal assistants: evidence from Croatia. *Information Research an international electronic journal*, 28 (2), 2–26.

Biloš, A., & **Budimir, B.** (2023). How much do we trust in AI? Exploring the impact of artificial intelligence on user trust levels. In *6. međunarodna naučna konferencija o digitalnoj ekonomiji DIEC 2023* (pp. 23-41). Tuzla: Visoka škola za savremeno poslovanje, informacione tehnologije i tržišne komunikacije "Internacionalna poslovno-informaciona akademija" Tuzla.

Biloš, A., **Budimir, B.**, & Zrilić, V. (2021, September). Influence of the digital environment on decorative cosmetics trends. In *27th CROMAR Congress Let the masks fall: New consumer in business and research (CROMAR 2021)* (pp. 21–43).

Biloš, A., **Budimir, B.**, & Vilk, A. (2024). Modno razotkrivanje: Primjena SUPR-Q modela za razumijevanje korisničkog iskustva na mrežnim sjedištima modnih marki u EU. In *9th International Scientific Conference CRODMA 2024: Book of Papers* (pp. 63-74). Varaždin: Croatian Academy of Sciences and Arts; University of Zagreb Faculty of Organization and Informatics; Croatian Direct Marketing Association.

Biloš, A., **Budimir, B.**, & Rašić, J. (2022). Rural areas' digital competitiveness: a comparative analysis. In *18th Interdisciplinary Management Research (IMR 2022)* (pp. 1052–1072).

Budimir, B., Kulić, K., & Pajcur, M. (2021). Mobile marketing as a component in consumer's behavior. In *17th Interdisciplinary Management Research (IMR 2021)* (Vol. 2, pp. 1167–1182).

Crnković, B., Rašić, J., & **Budimir, B.** (2020). The connection between investing in education and economic development. *IMR2020, interdisciplinary management research xvi, interdisziplinäre managementforschung*, (pp. 1687–1704).

Biloš, A., & **Budimir, B.** (2020). The Use of the Collaborative Economy in the EU: Senior Citizens Perspective. *Aging Society: Rethinking and Redesigning Retirement*, 55–78.