



Kauno technologijos universitetas

Ekonomikos ir verslo fakultetas

**Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose
tinkluose poveikio vartotojų pasitikėjimui bei ketinimams
lyginamoji analizė**

Baigiamasis magistro projektas

Airidas Kundrotas

Projekto autorius

Doc. Aistė Dovalienė

Vadovė

Kaunas, 2025



Kauno technologijos universitetas

Ekonomikos ir verslo fakultetas

**Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose
tinkluose poveikio vartotojų pasitikėjimui bei ketinimams
lyginamoji analizė**

Baigiamasis magistro projektas

Marketingo valdymas (6211LX038)

Airidas Kundrotas

Projekto autorius

Doc. Aistė Dovalienė

Vadovė

Doc. Agnė Gadeikienė

Recenzentė

Kaunas, 2025



Kauno technologijos universitetas

Ekonomikos ir verslo fakultetas

Airidas Kundrotas

Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams lyginamoji analizė

Akademinio sąžiningumo deklaracija

Patvirtinu, kad:

1. baigiamąjį projektą parengiau savarankiškai ir sąžiningai, nepažeisdama(s) kitų asmenų autoriaus ar kitų teisių, laikydamasi(s) Lietuvos Respublikos autorių teisių ir gretutinių teisių įstatymo nuostatų, Kauno technologijos universiteto (toliau – Universitetas) intelektinės nuosavybės valdymo ir perdavimo nuostatų bei Universiteto akademinės etikos kodekse nustatytų etikos reikalavimų;
2. baigiamajame projekte visi pateikti duomenys ir tyrimų rezultatai yra teisingi ir gauti teisėtai, nei viena šio projekto dalis nėra plagijuota nuo jokių spausdintinių ar elektroninių šaltinių, visos baigiamojo projekto tekste pateiktos citatos ir nuorodos yra nurodytos literatūros sąrašė;
3. įstatymų nenumatytų piniginių sumų už baigiamąjį projektą ar jo dalis niekam nesu mokėjęs (-usi);
4. suprantu, kad išaiškėjus nesąžiningumo ar kitų asmenų teisių pažeidimo faktui, man bus taikomos akademinės nuobaudos pagal Universitete galiojančią tvarką ir būsiu pašalinta(s) iš Universiteto, o baigiamasis projektas gali būti pateiktas Akademinės etikos ir procedūrų kontrolieriaus tarnybai nagrinėjant galimą akademinės etikos pažeidimą.

Airidas Kundrotas

Patvirtinta elektroniniu būdu

Kundrotas, Airidas. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams lyginamoji analizė. Magistro baigiamasis projektas / vadovė doc. Aistė Dovalienė; Kauno technologijos universitetas, Ekonomikos ir verslo fakultetas.

Studijų kryptis ir sritis (studijų krypčių grupė): Rinkodara, Verslas ir viešoji vadyba.

Reikšminiai žodžiai: dirbtinis intelektas, turinys socialiniuose tinkluose, žmogaus kuriamas turinys, vartotojų pasitikėjimas, vartotojų ketinimai, turinio autorystė.

Kaunas, 2025. 99 p.

Santrauka

Dirbtinio intelekto (DI) technologijų integracija į įvairias sritis, įskaitant mediciną, inžineriją, ir rinkodaros turinį socialiniuose tinkluose, reikšmingai keičia vartotojų kasdienybę. Pavyzdžiui, mažmeninės prekybos ir vartojimo prekių įmonės planuoja vidutiniškai skirti 3,32 proc. savo pajamų DI diegimui nuo 2025 metų. Šios investicijos apims tokias sritis kaip klientų aptarnavimas, tiekimo grandinės, talentų pritraukimas ir rinkodaros inovacijos, pabrėžiant DI plėtrą už tradicinių IT taikymo ribų (IBM, 2025a).

Nors DI naudojimas rinkodaroje dažnai vertinamas kaip būdas sumažinti kaštus ir padidinti efektyvumą (Ivanov, 2023), taip pat pastebimos reikšmingos problemos dėl vartotojų reakcijų į DI sukurtą turinį. Pavyzdžiui, tyrimai rodo, kad vartotojai dažnai nepastebi skirtumo tarp DI generuotų ir žmonių sukurtų reklamų, ypač kai naudojamos pažangios sistemos, tokios kaip „ChatGPT“ ar „Midjourney“. DI turinys dažnai pasiekia aukštą kokybės lygį, kuris gali sukelti vartotojams „technologinio svetimumo“ (angl. *perceived eeriness*) jausmą (Gu et al., 2024). Tyrimai taip pat atskleidžia, kad vartotojams dažnai reikia papildomos informacijos, jog būtų galima nustatyti, kas yra turinio autorius. Nepaisant to, DI sukurtas turinys dažnai laikomas objektyvesniu ir mažiau šališku nei žmogaus sukurtas, o tai gali padidinti socialiniuose tinkluose esančio turinio efektyvumą vartotojų tarpe tam tikruose kontekstuose (Xie et al., 2024). Socialiniai tinklai tampa ne tik platforma bendravimui, bet ir vieta, kur DI sprendimai gali optimizuoti turinio kūrimą, suteikti asmenines rekomendacijas bei analizuoti vartotojų elgseną realiu laiku (Mohamed et al., 2024). Vis dėlto šių technologijų naudojimas išlieka kontraversiškas, ypač kai vartotojai nesupranta, kaip DI grįstos sistemos priima vienus ar kitus sprendimus (Scharowski et al., 2023).

Baigiamajame magistro projekte siekiama palyginti palyginti dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikį vartotojų pasitikėjimui bei ketinimams. Šiam tikslui pasiekti formuojami penki tyrimo uždaviniai: pagrįsti dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams tyrimo aktualumą ir problematiką; atlikus teorinę dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams analizę, parengti conceptualų modelį; parengti į dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams atskleidimą orientuotą metodologiją; atlikus dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams empirinį tyrimą, interpretuoti rezultatus mokslinės diskusijos forma, išskirti tyrimo ribotumus bei tolimesnių tyrimų kryptis; pateikti praktines rekomendacijas, skirtas organizacijoms bei individualiems asmenims, siekiantiems pasitelkti ir efektyviai naudoti dirbtinio intelekto ir žmogaus kuriamą turinį socialiniuose tinkluose.

Atliktas empirinis tyrimas atskleidė, kad vartotojų suvokiamas turinio veržlumas ir įtaigumas daro teigiamą, o sintetiškumas – reikšmingai neigiamą poveikį pasitikėjimui turiniu. Nustatyta, kad dirbtinio intelekto kurto turinio autoriaus atskleidimas daro neigiamą įtaką pasitikėjimui turiniu, o žmogaus kurto turinio autoriaus atskleidimas – teigiamą. Tuo tarpu keistumas, priešingai nei prognozuota, turėjo teigiamą poveikį pasitikėjimui turiniu, suponuojantį, kad respondentams keistas turinys gali asocijuotis su inovatyvumu ir naujumu. Taip pat nustatyta, kad požiūris į dirbtinį intelektą kaip į bendraautorių moderuoja turinio autoriaus atskleidimo poveikį pasitikėjimui, tačiau požiūris į DI kaip į įrankį šio poveikio reikšmingai neveikia. Tyrimo duomenų analizė taip pat atskleidė, kad vartotojų pasitikėjimo pokytis turiniu teigiamai veikia jų polinkį sąveikauti su turiniu, tačiau šis ryšys buvo gana silpnas, nurodant, kad pasitikėjimas nėra vienintelis veiksnys, lemiantis vartotojų ketinimus. Taip pat pastebėta, kad nepriklausomai nuo respondentų subjektyvių žinių apie DI įrankius ar subjektyvaus gebėjimo atpažinti DI turinį, tinkamai turinio autorių identifikavo 32,5 proc. tyrimo dalyvių. Priešingai nei buvo prognozuota, tyrimo rezultatai atskleidė, kad respondentai, prieš atskleidžiant turinio autorių, DI turinį vertino kaip autentiškesnį, kokybiškesnį ir patikimesnį nei žmogaus sukurtą.

Kundrotas, Airidas. A Comparative Analysis of the Impact of Artificial Intelligence-Generated and Human-Generated Social Media Content on Users' Trust and Intentions. Master's Final Degree Project / supervisor Assoc. Prof. Aistė Dovalienė; School of Economics and Business, Kaunas University of Technology.

Study field and area (study field group): Marketing, Business and Public Management.

Keywords: artificial intelligence, content on social media, human-generated content, consumer trust, consumer intentions, content authorship.

Kaunas, 2025. 99 p.

Summary

The integration of artificial intelligence (AI) technologies into diverse sectors, including medicine, engineering, and social media marketing, is significantly transforming daily consumer experiences. For instance, retail and consumer goods companies plan to allocate an average of 3.32% of their revenues to AI implementation starting in 2025. These investments are expected to target areas such as customer service, supply chain management, talent acquisition, and marketing innovations, emphasizing the expansion of AI beyond traditional IT applications (IBM, 2025a).

While the use of AI in marketing is often seen as a way to reduce costs and improve efficiency (Ivanov, 2023), significant challenges related to consumer reactions to AI-generated content have also emerged. For example, studies indicate that consumers often struggle to distinguish between AI-generated and human-created advertisements, especially when using advanced systems like ChatGPT or Midjourney. AI-generated content frequently reaches a high level of quality, potentially triggering a sense of "technological eeriness" among users (Gu et al., 2024). Research also reveals that consumers often require additional context to identify the true author of the content. Nonetheless, AI-generated content is frequently perceived as more objective and less biased than human-created content, which can enhance the effectiveness of social media communications in certain contexts (Xie et al., 2024). Social media platforms are evolving into spaces not only for communication but also for AI-driven content creation, personalized recommendations, and real-time user behavior analysis (Mohamed et al., 2024). However, the use of these technologies remains controversial, particularly when users lack a clear understanding of how AI-powered systems make decisions (Scharowski et al., 2023).

The primary objective of this master's thesis is to compare the impact of AI-generated and human-created social media content on users' trust and intentions. To achieve this, the research is structured around five key goals: establish the relevance and significance of studying the impact of AI-generated and human-generated social media content on users' trust and intentions; develop a conceptual model based on a theoretical analysis of the effects of AI-generated and human-generated social media content on users' trust and intentions; design a methodology focused on revealing the impact of AI and human-generated social media content on users' trust and intentions within social media; conduct an empirical study to interpret the findings within a scientific discussion, highlighting the limitations of the research and directions for future studies; provide practical recommendations for organizations and individuals seeking to effectively utilize AI and human-generated social media content.

The empirical research revealed that perceived content vividness and persuasiveness have a positive impact on trust, while synthetic nature significantly negatively affects trust. It was also found that disclosing the authorship of AI-generated content negatively impacts trust, while revealing the authorship of human-created content has a positive effect. Contrary to expectations, perceived eeriness had a positive impact on trust, suggesting that respondents may associate unusual content with innovation and novelty. Additionally, the study found that the perception of AI as a co-author moderates the impact of author disclosure on trust, while the perception of AI as merely a tool does not significantly influence this relationship. The data analysis also indicated that changes in consumer trust positively influence their inclination to interact with the content, though this relationship was relatively weak, suggesting that trust is not the only factor driving consumer intentions. Moreover, regardless of subjective knowledge about AI tools or perceived ability to identify AI-generated content, only 32.5% of participants accurately identified the content author. Surprisingly, before the author was disclosed, respondents rated AI-generated content as more authentic, higher quality, and more trustworthy than human-created content, contrary to initial expectations.

Turinys

Turinys.....	7
Lentelių sąrašas	9
Paveikslų sąrašas	11
Įvadas.....	12
1. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams tyrimo aktualumas ir problematika	14
2. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams tyrimo teorinis pagrindimas.....	25
2.1. DI ir žmogaus kuriamo turinio socialiniuose tinkluose charakteristikos	25
2.1.1. Socialiniuose tinkluose kuriamo komunikacinio turinio tipai bei savybės	25
2.1.2. DI panaudojimas socialiniuose tinkluose bei jo pagalba sukurtu turinio ypatumai	27
2.1.3. Žmogaus kuriamo turinio socialiniuose tinkluose charakteristikos	31
2.2. Vartotojų pasitikėjimas dirbtiniu intelektu ir jo etiniai aspektai	34
2.3. Turinio socialiniuose tinkluose autentiškumo suvokimas.....	38
2.4. Vartotojų pasitikėjimas bei ketinimai dirbtinio intelekto ir žmogaus kuriamo turinio atveju .	41
2.5. Dirbtinio intelekto ir žmogaus kuriamo turinio poveikio vartotojų pasitikėjimui ir ketinimams konceptualaus modelio pagrindimas	45
3. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams empirinio tyrimo metodologija	50
3.1. Empirinio tyrimo tikslas, uždaviniai ir hipotezės.....	50
3.2. Empirinio tyrimo tipas, duomenų rinkimo metodas ir tyrimo konstrukto operacionalus apibūdinimas.....	52
3.3. Empirinio tyrimo imties atrankos procedūros, tyrimo etika ir analizės metodai	57
4. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams empirinį tyrimo rezultatai	59
4.1. Bendri respondentų demografiniai ir sociografiniai rodikliai	59
4.2. Tyrimo instrumento patikimumo ir kokybės įvertinimas	61
4.3. Psichometriniais matavimais vertinamų kintamųjų charakteristikos	66
4.4. Kintamųjų tarpusavio ryšių analizė ir hipotezių tikrinimas	71
4.5. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose įtakos vartotojų pasitikėjimui bei ketinimams tyrimo rezultatų apibendrinimas, mokslinė diskusija, tyrimo ribotumai bei tolimesnių tyrimų kryptys.....	83
Tyrimo išvados ir pasiūlymai	89
Literatūros sąrašas	92
Informacijos šaltinių sąrašas	98
Priedai.....	100
1 priedas. Pirminės tyrimo hipotezės ir jų detalizacija.....	100
2 priedas. Tyrimo anketa	102
3 priedas. „ChatGPT“ ir „Gemini“ užklausa	107
4 priedas. Įrašas apie anketą socialiniuose tinkluose.....	108
5 priedas. Respondentų demografiniai duomenys.....	109
6 priedas. Tyrimo konstrukto ir skalių kokybės vertinimas	113
7 priedas. Psichometriniai matavimai.....	118

8	priedas. Koreliacinė kintamųjų tarpusavio ryšių analizė.....	122
9	priedas. Regresinė analizė	123
10	priedas. Neparametrinis Mann-Whitney U testas.....	123
11	priedas. Paprastoji regresinė analizė.....	124
12	priedas. Moderacijos analizė	126
13	priedas. Chi-kvadrato testas.....	126

Lentelių sąrašas

1 lentelė. Pasitikėjimo dirbtiniu intelektu kuriamu turiniu mokslinių tyrimų apžvalga	15
2 lentelė. Pasitikėjimo dirbtiniu intelektu poveikis vartotojų ketinimams mokslinių tyrimų apžvalga	17
3 lentelė. Pasitikėjimas žmogaus sukurtu turiniu socialiniuose tinkluose	19
4 lentelė. Pagrindinės tematikos moksliniuose darbuose apie pasitikėjimą DI ir žmogaus kuriamu turiniu	22
5 lentelė. Populiariausi socialiniai tinklai Lietuvoje bei jų turinio formatų charakteristikos	26
6 lentelė. Turino kokybės vertinimas	27
7 lentelė. Dirbtinio intelekto socialiniuose tinkluose sritys	29
8 lentelė. Užklauso generatyviniam dirbtiniam intelektui struktūra ir pavyzdys (sukurta autoriaus remiantis IBM, 2025b)	29
9 lentelė. Žmogaus kuriamo turinio charakteristikos	31
10 lentelė. Atsakomybė už DI klaidas pagal subjektą (sudaryta remiantis ES dirbtinio intelekto aktu)	36
11 lentelė. Vartotojų elgsena DI ir žmogaus kuriamo turinio atveju	43
12 lentelė. Vartotojų pasitikėjimas DI ir žmogaus kuriamu turiniu	44
13 lentelė. Tyrimo konstruktai ir matavimo skalės	55
14 lentelė. Imties dydis moksliniuose darbuose	57
15 lentelė. Respondentų lyties ir išsilavinimo pasiskirstymas (N=461)	59
16 lentelė. Respondentų pasiskirstymas pagal amžių (N=461)	60
17 lentelė. Respondentų socialinių tinklų naudojimo tendencijos (N=461)	60
18 lentelė. Skalių tinkamumo vertinimas (N=461)	61
19 lentelė. Veiksnių faktorinė analizė (N=461)	62
20 lentelė. Atnaujintos skalės, priskiriamos veiksniams (N=461)	64
21 lentelė. Pasekmių faktorinė analizė	64
22 lentelė. Tyrimo skalių patikimumo vertinimas (N=461)	65
23 lentelė. Atnaujintų hipotezių sąrašas	65
24 lentelė. Psichometriniais matavimais vertinamų kintamųjų charakteristikos (N=461)	67
25 lentelė. Kintamųjų reikšmių priklausomybė nuo lyties (N=460)	68
26 lentelė. Kolmogorov'o-Smirnov'o (K-S) normalumo testo rezultatai (N=461)	69
27 lentelė. Kintamųjų reikšmių priklausomybė nuo amžiaus, išsilavinimo, socialinių tinklų naudojimo dažnio, trukmės bei DI įrankių naudojimo dažnio (N=461)	70
28 lentelė. Kintamųjų kodavimas koreliacinei analizei	71
29 lentelė. Koreliacijos koeficientų interpretavimas pagal Kafle (2023)	72
30 lentelė. Koreliacinės analizės tarp tyrimui svarbių kintamųjų rezultatai (N=461)	72
31 lentelė. Daugialypė regresija tarp autentiškumo ir kokybės bei pasitikėjimo. Modelio tinkamumas (N=461)	74
32 lentelė. Daugialypė regresija tarp autentiškumo ir kokybės bei pasitikėjimo. Nepriklausomų kintamųjų įtaka (N=461)	74
33 lentelė. Neparametrinis Mann-Whitney U testas tarp turinio autoriaus atskleidimo ir pasitikėjimo pokyčio (N=461)	75
34 lentelė. Paprastoji regresija tarp pasitikėjimo pokyčio ir elgsenos. Modelio tinkamumas (N=461)	75

35 lentelė. Paprastoji regresija tarp pasitikėjimo pokyčio ir elgsenos. Nepriklausomojo kintamojo įtaka (N=461).....	76
36 lentelė. Moderacijos analizė tarp autorystės atskleidimo, pasitikėjimo pokyčio ir požiūrio į DI. Modelio tinkamumas (N=461)	76
37 lentelė. Moderacijos analizė tarp autorystės atskleidimo, pasitikėjimo pokyčio ir požiūrio į DI. Rezultatai (N=461)	76
38 lentelė. Chi-kvadrato analizė tarp DI žinių kategorijų ir autoriaus identifikavimo (N=461)	77
39 lentelė. Chi-kvadrato analizė tarp subjektyvaus gebėjimo identifikuoti turinio autorių ir autoriaus identifikavimo (N=461).....	79
40 lentelė. Paprastoji regresija tarp suvokiamo pasitikėjimo didėjimo sužinojus autorių ir faktinio pasitikėjimo pokyčio. Modelio tinkamumas (N=461).....	80
41 lentelė. Neparametrinis Mann-Whitney U testas tarp turinio autoriaus atskleidimo bei autentiškumo ir kokybės dimensijų (N=461)	80
42 lentelė. Neparametrinis Mann-Whitney U testas tarp turinio autoriaus bei pasitikėjimo turiniu (N=461)	81
43 lentelė. Tyrimo hipotezių rezultatai	82

Paveikslų sąrašas

1 pav. Etinės DI problemos (pagal Stahl et al., 2021).....	35
2 pav. „Midjourney“ sukurta vaizdo klastotė (autorius nežinomas).....	42
3 pav. Konceptualus dirbtinio intelekto ir žmogaus kuriamo turinio poveikio vartotojų pasitikėjimui ir ketinimams modelis	48
4 pav. Tyrimo modelis	50
5 pav. Anketos schema.....	53
6 pav. Respondentų subjektyvių žinių apie DI įrankius ir faktinio autoriaus identifikavimo kryžminė analizė (N=461).....	78
7 pav. Respondentų suvokiamo gebėjimo atskirti DI turinį nuo žmogaus kurto turinio ir faktinio autoriaus identifikavimo kryžminė analizė (N=461).....	79

Ivadas

Tyrimo aktualumas. Dirbtinio intelekto (DI) technologijų integracija į įvairias sritis, įskaitant mediciną, inžineriją ir rinkodaros turinį socialiniuose tinkluose, reikšmingai keičia vartotojų kasdienybę. Pavyzdžiui, mažmeninės prekybos ir vartojimo prekių įmonės planuoja vidutiniškai skirti 3,32 proc. savo pajamų DI diegimui nuo 2025 metų. Šios investicijos apims tokias sritis kaip klientų aptarnavimas, tiekimo grandinės, talentų pritraukimas ir rinkodaros inovacijos, pabrėžiant DI plėtrą už tradicinių IT taikymo ribų (IBM, 2025a).

Socialinių tinklų platformose DI algoritmai kuria turinį, kuris yra vizualiai patrauklus ir dažnai vertinamas kaip aukštos kokybės. Statistiniai duomenys rodo, kad DI sukurtas turinio poveikis sparčiai auga: 2023-iaisiais „AIART“ žyma „TikTok“ platformoje pasiekė 7,4 mlrd. peržiūrų, o „Instagram“ su DI susijusios žymos buvo panaudotos daugiau nei 19 milijonų kartų (Park et al., 2024). DI algoritmai šiuo metu generuoja ir optimizuoja turinį, kuris siekia atliepti įvairius vartotojų lūkesčius bei kurti suasmenintas patirtis (Saheb et al., 2024). Tyrimai rodo, kad DI grindžiami sprendimai socialinių tinklų platformose padidina vartotojų įsitraukimą, nes leidžia pateikti suasmenintą ir aktualų turinį, tačiau tuo pačiu metu kyla iššūkių, susijusių su etiniais klausimais, tokiais kaip skaidrumo bei patikimumo užtikrinimas (Scharowski et al., 2023). 2023 m. net 70 proc. rinkodaros specialistų naudojo DI technologijas siekdami suasmeninti turinį ir pagerinti vartotojų patirtį. Tokios technologijos leido padidinti rinkodaros kampanijų efektyvumą 2,5 karto ir sutrumpinti turinio kūrimo laiką iki 40 proc. (IBM, 2023).

Nors DI naudojimas rinkodaroje dažnai pasitelkiamas, siekiant sumažinti kaštus ir padidinti efektyvumą (Ivanov, 2023), taip pat pastebimos reikšmingos problemos dėl vartotojų reakcijų į DI sukurtą turinį. Pavyzdžiui, tyrimai rodo, kad vartotojai dažnai nepastebi skirtumo tarp DI generuotų ir žmonių sukurtų reklamų, ypač kai naudojamos pažangios sistemos, tokios kaip „ChatGPT“ ar „Midjourney“. DI turinys dažnai pasiekia aukštą kokybės lygį, kuris gali sukelti vartotojams „technologinį svetimumo“ (angl. *perceived eeriness*) jausmą (Gu et al., 2024). Tyrimai taip pat atskleidžia, kad vartotojams dažnai reikia papildomos informacijos, jog būtų galima nustatyti, kas yra turinio autorius. Nepaisant to, DI sukurtas turinys dažnai laikomas objektyvesniu ir mažiau šališku nei žmogaus sukurtas, o tai gali padidinti socialiniuose tinkluose esančio turinio efektyvumą vartotojų tarpe tam tikruose kontekstuose (Xie et al., 2024). Socialiniai tinklai tampa ne tik platforma bendravimui, bet ir vieta, kur DI sprendimai gali optimizuoti turinio kūrimą, suteikti asmenines rekomendacijas bei analizuoti vartotojų elgseną realiu laiku (Mohamed et al., 2024). Vis dėlto šių technologijų naudojimas išlieka kontraversiškas, ypač kai vartotojai nesupranta, kaip DI grįstos sistemos priima vienus ar kitus sprendimus (Scharowski et al., 2023).

Tyrimo problema. Analizuojant akademinis šaltinius, orientuotus į žmogaus ir DI kurto turinio sąveiką su vartotojais, pastebima, kad vis dar trūksta tyrimų, nagrinėjančių vartotojų požiūrį į dirbtinio intelekto ir žmogaus kuriamą tekstinį bei vizualinį turinį socialiniuose tinkluose. Nors kai kurie akademiniai darbai jau tiria pasitikėjimo aspektus dirbtinio intelekto ir žmogaus generuojamo turinio kontekste, vis dar stokojama išsamios analizės apie tai, kaip skirtingos turinio savybės, tokios kaip autentiškumas, autorystė ir šaltinio patikimumas, formuoja vartotojų elgseną ir sprendimus (Gu et al., 2024; Larsson & Heintz, 2020; E. A. S. Mohamed et al., 2024a; Park et al., 2024; Xie et al., 2024). Pasitikėjimas yra kertinis veiksnys, užtikrinantis sėkmingą žmogaus ir DI sąveiką (Ferrario et al., 2020; Ulfert et al., 2024). Praktiškai tai reiškia, kad tyrimai, nagrinėjantys pasitikėjimą DI kontekste, gali padėti stiprinti komandų efektyvumą, mažinti klaidų tikimybę ir skatinti DI pritaikymą

įvairiose organizacijose (Ulfert et al., 2024). Be to, būtina atsižvelgti į tai, kaip šios savybės veikia vartotojų emocinį įsitraukimą ir ilgalaikį lojalumą platformoms ar prekių ženklams (Y. Chen et al., 2022). Šis mokslinių žinių trūkumas išryškina poreikį giliau nagrinėti šiuos aspektus, siekiant pateikti išsamesnę vartotojų elgsenos analizę, apimančią tiek tekstinį, tiek vizualinį turinį (Park et al., 2024). Taip pat svarbu atkreipti dėmesį į dirbtinio intelekto generuojamo turinio skaidrumą ir paaiškinamumą, nes šie aspektai gali daryti reikšmingą įtaką vartotojų pasitikėjimui ir jų norui naudotis tokiomis technologijomis (Felzmann et al., 2020; Larsson & Heintz, 2020; Wang, 2022). Tokie tyrimai ne tik papildytų mokslinę literatūrą, bet ir suteiktų vertingų įžvalgų turinio kūrėjams bei socialinių tinklų platformų administratoriams apie tai, kaip stiprinti vartotojų pasitikėjimą, skatinti teigiamus sprendimus ir kurti etišką bei į žmogų orientuotą socialinių tinklų ekosistemą. Tai ypač svarbu nagrinėjant, kaip turinio autentiškumas, autorystės atskleidimas ir technologinės galimybės veikia vartotojų pasitikėjimą bei ketinimus skaitmeninėje erdvėje.

Remiantis aukščiau išvardintais argumentais, formuojama **mokslinė problema**: kaip dirbtinio intelekto ir žmogaus kuriamas turinys socialiniuose tinkluose veikia vartotojų pasitikėjimą bei ketinimus?

Objektas: dirbtinio intelekto ir žmogaus kuriamas turinys socialiniuose tinkluose ir jo poveikis vartotojų pasitikėjimui bei ketinimams.

Tikslas: teoriškai ir empiriškai palyginti dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikį vartotojų pasitikėjimui bei ketinimams.

Uždaviniai:

1. pagrįsti dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams tyrimo aktualumą ir problematiką;
2. atlikus teorinę dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams analizę, parengti konceptualų modelį;
3. parengti į dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams atskleidimą orientuotą metodologiją;
4. atlikus dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams empirinį tyrimą, interpretuoti rezultatus mokslinės diskusijos forma, išskirti tyrimo ribotumus bei tolimesnių tyrimų kryptis;
5. pateikti praktines rekomendacijas, skirtas organizacijoms bei individualiems asmenims, siekiantiems pasitelkti ir efektyviai naudoti dirbtinio intelekto ir žmogaus kuriamą turinį socialiniuose tinkluose.

Tyrimo metodai: mokslinės literatūros analizė; kiekybinis scenarijais grįstas eksperimentinis empirinis tyrimas; duomenų rinkimas internetine anketa; duomenų analizė naudojant „IBM SPSS“; aprašomoji analizė; faktorinė analizė; koreliacinė analizė; tiesinės ir daugialypės regresijos analizė; moderavimo analizė; Mann-Whitney U testas; Chi-kvadrato testas.

1. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams tyrimo aktualumas ir problematika

Dirbtinio intelekto sprendimams apimant vis daugiau sričių, kūrybinio turinio kūrimas ir suvokimas fundamentaliai keičiasi. Dirbtinio intelekto gebėjimas sukurti aukštos kokybės meninį, muzikinį, literatūrinį ir vizualinį turinį ištrynė ribas tarp žmogaus ir DI kūrybiškumo. Įvairių socialinių medijų platformose, dirbtiniu intelektu paremti įrankiai kuria viską: nuo tekstines ar vizualinės informacijos, iki muzikos ar visiškai automatizuotų įrašų (Park et al., 2024). Tuo tarpu socialinės medijos yra laikomos kaip vienas iš esminių įrankių šiuolaikinėje komunikacijoje, atliekančiu esminį vaidmenį formuojant visuomenės nuomonę, palaikant tarpasmeninius santykius, ar kuriant ryšį tarp verslo ir jo kliento (Saheb et al., 2024).

Šiame skyriuje apžvelgiamas ir aptariamas mokslinės literatūros, susijusios su baigiamojo magistro projekto tema, ištirtumo lygis. Atkreipiant dėmesį į tai, jog darbo tikslas – išsiaiškinti, kokį poveikį dirbtinio intelekto ir žmogaus kuriamas turinys daro vartotojų pasitikėjimui bei ketinimams socialiniuose tinkluose, orientuojamasi į dirbtinio intelekto technologijų ir žmogaus kuriamo turinio palyginimą. Siekiant tinkamai atskleisti tyrimo aktualumą ir problematiką, identifikuotos dvi pagrindinės tyrimų tematikos: pasitikėjimas dirbtinio intelekto kuriamu turiniu bei pasitikėjimas žmogaus kuriamu turiniu.

Pasitikėjimo dirbtiniu intelektu ištirtumo lygmuo yra ganėtinai aukštas, nagrinėjantis daug skirtingų pasitikėjimo DI aspektų. Visgi, atkreipiant dėmesį į skirtingas metodologijas ir atlikimo kontekstus, tyrimų rezultatai neretai skiriasi. Toliau pateikiamas keleto tyrimų apie pasitikėjimą dirbtinio intelekto technologijomis apibendrinimas, atskleidžiantis skirtingų mokslinių darbų išvadas.

Scharowski'o et al. (2023) tyrimas nagrinėja skaidrumą per dvi į žmogų orientuotas DI paaiškinimo rūšis: vienas sutelktas į paaiškinimą, kokie veiksniai turėjo įtakos vienam ar kitam DI sprendimui, tuo tarpu kita rūšis orientuota į alternatyvius „o jeigu“ scenarijus, kurie paaiškintų kaip vieni ar kiti pradiniai įvesties duomenys būtų pakeitę rezultatą. Tyrejai siekė atsakyti į klausimą kaip į žmogų orientuotas DI įrankių paaiškinimas galėtų padidinti pasiklovimą DI rekomendacijomis. Tyrimo rezultatai atskleidė, jog į žmogų orientuoti paaiškinimai tam tikromis sąlygomis padidino pasiklovimą DI siūlomomis rekomendacijomis. Anot Scharowski'o et al. (2023), DI įrankių veikimo principų paaiškinimai naudotojų nuomonės, t. y. pasitikėjimo neveikia tiek, kiek veikia jų elgesį, t. y. pasiklovimą, taip pabrėžiant, jog svarbu atskirti pasitikėjimą (nuomonę) ir pasiklovimą (elgesį), siekiant tiksliai atsispindėti vartotojų sąveiką su DI įrankiais.

Tyrimas, kurį atliko Y. Chen et al. (2022) atskleidė, jog prielaida, kad DI integracija į verslo veiklą savaime gerina vartotojo patirtį, tiesiogiai priklauso nuo konteksto, t. y. verslo šakos, į kurią DI sprendimai yra integruojami. Pavyzdžiui, DI integracija į trumpalaikės nuomos paslaugas teikiančias platformas, tokias kaip „Airbnb“ daro neigiamą įtaką vartotojo pasitikėjimui, įsitraukimui ir lojalumui, nes pasigendama žmogiškosios sąveikos. Itin daug DI sistemų nėra plačiai pripažįstamos dalinai dėl to, jog vartotojai jų pilnai nesupranta arba nepasitiki sprendimų priėmimo procesu.

Anot Larsson'o ir Heintz'o (2020), išanalizavus skaidrumo vaidmenį dirbtinio intelekto sprendimų priėmimo kontekste, pabrėžiama, jog vien techninių paaiškinimų nepakanka – būtina įtraukti ir teisinių, socialinių, bei humanitarinių sričių ekspertizę. Akcentuojama, jog siekiant užtikrinti skaidrumą, būtina atsižvelgti į teisinius DI aspektus, į žmonių duomenų suvokimo lygį, užkirsti kelią DI įrankių išnaudojimui, aiškiai iškomunikuoti DI veikimo principus, atsižvelgti į jautrios

informacijos apsaugos sistemas bei personalizuotą turinį. Multidisciplininis požiūris į DI skaidrumą yra būtinas, siekiant kurti pasitikėjimą vartotojų tarpe.

Ferrario et al. (2020) savo tyrime pristatė „inkrementinį pasitikėjimo modelį“, skirtą analizuoti žmogaus ir DI sąveikas. Tyrimas koncentruotas į pasitikėjimo dinamiką žmogaus ir DI sąveikose, išskiriant kognityvinius ir nekognityvinius aspektus. Mokslininkų pasiūlytas modelis, kuris įvertinų pasitikėjimo lygius – paprastą, refleksyvų ir paradigminį pasitikėjimą – modeliu paremti rezultatai būtų taikytini įvairiems verslo ir socialiniams scenarijams. Tyrimo rezultatai parodė, kad pasitikėjimo lygis priklauso nuo episteminių, pavyzdžiui, DI tikslumo, ir pragmatinių, pavyzdžiui, sąnaudų taupymo, priežasčių. Tuo tarpu paradigminis pasitikėjimas, kaip pilnavertė pasitikėjimo forma, reikalauja tiek pasiklovimo, tiek tikėjimo DI patikimumu.

Ulfert et al. (2024) savo tyrime nagrinėjo pasitikėjimo formavimosi procesus komandose, sudarytose iš žmonių ir dirbtinio intelekto sistemų. Mokslininkai siekė sukurti daugiadisciplinę teorinę struktūrą, kuri apibūdintų, kaip pasitikėjimas vystosi skirtingais lygmenimis – individualiu, poros ir komandos. Tyrimas atskleidė, kad pasitikėjimas DI sistemomis priklauso nuo jų gebėjimo pagrįsti vienus ar kitus sprendimus bei demonstruoti patikimumą, kuris apima gebėjimus, integralumą ir geranoriškumą. Be to, nustatyta, kad DI agentai gali veikti ne tik kaip pasitikėjimo objektai, bet ir kaip aktyvūs komandos nariai. Galiausiai nustatyta, jog tinkamai integruotos DI sistemos gali didinti komandos darbo efektyvumą, tačiau būtina spręsti nepasitikėjimo arba per didelio pasitikėjimo jais problemas.

Yang ir Rau'as (2024) nagrinėjo pasitikėjimo formavimąsi žmogaus ir DI partnerių sąveikose, lyginant investavimo elgseną, emocinį sužadinimą ir neuroveiklą. Mokslininkai analizavo pasitikėjimo kūrimo procesą dviem etapais: pradinį pasitikėjimą ir pasitikėjimo kitimą pasitelkiant tiriamųjų subjektų grįžtamąjį ryšį. Rezultatai parodė, kad pradinis pasitikėjimas DI turėjo didesnę įtaką galutiniam pasitikėjimui nei pradinis pasitikėjimas žmogumi, tačiau žmogaus grįžtamasis ryšys stipriau paveikė pasitikėjimo vystymą ir emocinį sužadinimą. Dalyviai jautė daugiau teigiamų emocijų ir stipresnę prefrontalinės žievės, atsakingos už svarbias kognityvines ir emocines funkcijas, aktyvaciją bendraudami su žmonėmis nei su DI. Tyrimo rezultatai pabrėžė žmogaus atgalinio ryšio svarbą pasitikėjimo kūrimui, tuo tarpu pasitikėjimas dirbtiniu intelektu labiau priklauso nuo pradinio įvertinimo.

Analizuojant pasitikėjimo tematiką DI kontekste, pastebima, jog dalis daugiau dėmesio skiria pačiam pasitikėjimui, tuo tarpu kiti orientuojasi į pasitikėjimo įtaką ketinimams. Žemiau pateikiamoje lentelėje (žr. 1 lentelę), atsispindi koncentruota tyrimų apie pasitikėjimą dirbtiniu intelektu apžvalga, apimanti tyrimo objektą, tikslus, hipotezes bei rezultatus.

1 lentelė. Pasitikėjimo dirbtiniu intelektu kuriamu turiniu mokslinių tyrimų apžvalga

Autorius(-iai), metai	Tyrimo objektas	Tyrimo tikslai / hipotezės	Tyrimo rezultatai
Larsson ir Heintz, (2020).	Dirbtinio intelekto skaidrumas.	Išnagrinėti DI skaidrumo sampratą iš socio-teisinės ir kompiuterių mokslų perspektyvų. Hipotezė: skaidrumas yra daugialypis, o DI skaidrumui reikalingas tarpdisciplininis požiūris, apimantis tiek teises, tiek socialines normas bei DI valdymą.	Skaidrumas DI yra daugialypis, apimantis algoritminį skaidrumą ir platesnį DI skaidrumą. DI skaidrumas turėtų būti aiškesnis vartotojams, įtraukiant tiek teisinius, tiek socialinius aspektus. Siūloma plėtoti „skaidrumo“ sampratą kaip platesnę nei vien tik algoritminį aspektą, kad būtų sukurta patikima DI valdymo sistema, prisidedanti

Autorius(-iai), metai	Tyrimo objektas	Tyrimo tikslai / hipotezės	Tyrimo rezultatai
			prie etinio ir socialiai atsakingo DI vystymo.
Ferrario et al. (2020).	Pasitikėjimas žmogaus ir DI sąveikoje.	Išanalizuoti pasitikėjimo dinamiką žmogaus ir DI sąveikoje, sukuriant daugiasluoksnį modelį, kuris jungia kognityvinius ir nekognityvinius pasitikėjimo aspektus. Hipotezės: 1) pasitikėjimas DI priklauso nuo episteminių (žinojimo) ir pragmatinių (praktinių) priežasčių; 2) refleksyvus ir paradigminis pasitikėjimas yra svarbiausi pilnavertės pasitikėjimo formos elementai; 3) pasitikėjimas DI gali būti santykinis (priklausomai nuo aplinkybių) arba absoliutus.	Pristatomas „inkrementinis pasitikėjimo modelis“, apimantis tris sluoksnius: 1) Paprastas pasitikėjimas. Kai žmogus pasikliauja DI be kontrolės arba patikimumo analizės; 2) Refleksyvus pasitikėjimas. Paremtas tikėjimu DI, dažnai pagrįstu objektyviais duomenimis; 3) Paradigminis pasitikėjimas. Kai paprastas ir refleksyvus pasitikėjimas egzistuoja kartu. Modelis taikomas situacijoms, kai DI naudojamas versle, pavyzdžiui, personalizuotoms prognozėms generuoti. Straipsnyje išskiriami episteminiai, t. y. objektyvūs ir subjektyvūs, bei pragmatiniai, pavyzdžiui, sąnaudų mažinimo, pasitikėjimo pagrindai. Galiausiai, autoriai pabrėžia, kad absoliutus pasitikėjimas DI yra idealas, kurį sunku pasiekti dėl DI dinamiškumo ir jo apmokymo procesų.
Ulfert (2024)	Pasitikėjimo dinamika žmogaus ir DI komandose, apimanti individualų, poros ir komandinį sąveikos lygmenis.	Sukurti daugiadisciplinę ir daugialygę teorinę struktūrą, apibrėžiančią pasitikėjimą komandose, kuriose dirba žmonės ir DI. Išnagrinėti, kaip formuojasi pasitikėjimo santykiai tarp žmonių ir DI agentų skirtinguose lygmenyse: individualiame, poros ir komandos lygyje. Hipotezės: 1) pasitikėjimas tarp žmonių ir DI sistemų yra dvikryptis (žmogus ir DI sistema gali būti ir pasitikėjimo objektu, ir subjektu). 2) komandinį pasitikėjimą formuoja įvairių santykių (individualių, porinių, komandinių) pasitikėjimo lygis.	Tyrimas išryškino tris pasitikėjimo lygmenis: pasitikėjimą atskirais komandos nariais, t. y. žmonėmis ar DI, poros santykius, t. y. žmogus-žmogus, žmogus-DI, DI-DI, ir komandos kaip visumos pasitikėjimą. Išsiaiškinta, jog pasitikėjimas formuojamas abipusiai, kai kiekvienas narys vertina kitų narių patikimumą, o tai reiškia, jog ne tik žmogus vertina DI patikimumą, bet ir DI vertina žmogaus ar kitų DI sistemų patikimumą, pasitelkdamas specifinius algoritmus ar modelius. Žmonėms lengviau įvertinti kitų žmonių patikimumą nei DI, o DI pasitikėjimą labiau veikia jų techniniai gebėjimai ir skaidrumas.
Yang ir Rau (2023)	Pasitikėjimo formavimas tarp žmogaus ir DI, lyginant su žmogaus-žmogaus sąveika investicijų, emocijų ir neuroveiklos kontekstuose.	Ištirti pasitikėjimo formavimą dviem etapais: pradinį pasitikėjimą ir pasitikėjimo kitimą pasitelkiant grįžtamąjį ryšį, lyginant žmogaus ir DI subjektų įtaką. Hipotezės: 1) pradinis pasitikėjimas DI stipriau veikia galutinį pasitikėjimą nei pradinis pasitikėjimas žmogumi;	Pradinis pasitikėjimas DI turėjo didesnę įtaką galutiniam pasitikėjimui nei pradinis pasitikėjimas žmogumi. Tuo tarpu grįžtamasis ryšys iš žmogaus turėjo stipresnę įtaką pasitikėjimo formavimui ir emociniam sužaditimui. Dalyviai jautė daugiau teigiamų emocijų bendraudami su žmogumi nei su DI. Neuroveiklos duomenys parodė stipresnę prefrontalinės žievės, atsakingos už

Autorius(-iai), metai	Tyrimo objektas	Tyrimo tikslai / hipotezės	Tyrimo rezultatai
		2) žmogaus grįžtamasis ryšys stipriau veikia pasitikėjimą nei DI grįžtamasis ryšys; 3) sąveika su žmogumi sukelia stipresnes emocijas ir neuroveiklos reakcijas nei su DI.	svarbias kognityvines ir emocines funkcijas, veiklą sąveikoje su žmogumi, ypač sprendimų priėmimo metu. Rezultatai parodė, kad grįžtamasis ryšys iš žmogaus yra svarbesnis pasitikėjimo kūrimui nei iš DI.

Toliau (žr. 2 lentelę), pateikiama susisteminta informacija apie tyrimus, siekiančius išanalizuoti kaip pasitikėjimas DI veikia vartotojų požiūrį bei ketinimus.

2 lentelė. Pasitikėjimo dirbtiniu intelektu poveikis vartotojų ketinimams mokslinių tyrimų apžvalga

Autorius(-iai), metai	Tyrimo objektas	Tyrimo tikslai / hipotezės	Tyrimo rezultatai
Scharowski et al. (2023).	Į žmogų orientuoti DI paaiškinimai ir jų poveikis.	Išanalizuoti, kaip į žmogų orientuoti paaiškinimai veikia vartotojų pasitikėjimą ir pasiklovimą DI sistemomis. Hipotezės: 1) eksperimentinių grupių pasitikėjimas ir pasiklovimas DI sistemomis bus didesnis dėl funkcijų svarbos paaiškinimų, lyginant su kontroline grupe, kurios dalyviams paaiškinimai nebus pateikiami; 2) kontrafaktiniai (angl. <i>what if</i>) paaiškinimai bus efektyvesni nei už funkcijų svarbos paaiškinimus; 3) pasiklovimo lygis bus didesnis už pasitikėjimo lygį.	Nustatyta, jog į žmogų orientuoti paaiškinimai reikšmingai padidina pasiklovimą DI tik tam tikrais atvejais, pavyzdžiui, kai DI rekomenduoja mažinti kainą, tačiau neturi reikšmingo poveikio pasitikėjimui. Paaiškėjo, jog vartotojai labiau linkę pasikloviti rekomendacijomis mažinti kainą – reiškinys, kuris gali būti susijęs su žmogaus įgytu šališkumu minimizuoti nuostolius. Atskleista, kad pasiklovimas nėra tiesiogiai susijęs su pasitikėjimu, o tai pabrėžia būtinybę atskirti šias dvi sąvokas ir naudoti skirtingus vertinimo metodus jų analizei.
Y. Chen et al. (2022).	Klientų pasitikėjimas, DI poveikis klientų įsitraukimui ir lojalumui dalijimosi ekonomikos kontekste.	Ištirti, kaip pasitikėjimas namų dalijimosi platformų šeimininkais ir platformomis veikia klientų įsitraukimą ir lojalumą bei DI vaidmuo šioje sąveikoje. Hipotezės: 1) pasitikėjimas platforma teigiamai veikia pasitikėjimą šeimininkais; 2) pasitikėjimas šeimininkais ir platforma didina klientų įsitraukimą ir lojalumą; 4) DI vaidmuo gali būti neigiamas moderatorius dalijimosi platformų, šeimininkų ir klientų sąveikoje.	Pasitikėjimas šeimininkais ir platforma teigiamai veikia klientų įsitraukimą ir lojalumą. DI, kaip moderatorius, turėjo neigiamą poveikį – jis mažino klientų, pasitikinčių platforma ir šeimininkais, įsitraukimą ir lojalumą. Neapčiuopiamas DI pobūdis ir žmonių teikiama pirmenybė tiesioginiam kontaktui gali būti priežastys, kodėl DI neigiamai paveikė santykius tarp vartotojų ir namų dalijimosi platformos arba tarp vartotojų ir šeimininkų.

Remiantis įvairių mokslininkų darbų analize, tyrimai apie pasitikėjimą dirbtiniu intelektu atskleidžia skirtingas išvadas, priklausomai nuo konteksto ir analizuojamų veiksnių. Viena grupė autorių akcentuoja teigiamą į žmogų orientuotų paaiškinimų ir paradigminio pasitikėjimo poveikį, kuris gali skatinti pasiklovimą DI sprendimais. Pavyzdžiui, Scharowski's et al. (2023) nustatė, kad funkcijų svarbos ir kontrafaktiniai paaiškinimai tam tikrais atvejais padidina pasiklovimą DI sistemomis, tačiau neturi reikšmingo poveikio pasitikėjimui. Kita vertus, Ferrario et al. (2020) pristatė

„inkrementinio pasitikėjimo modelį“, kuris išryškina paradigminio pasitikėjimo svarbą, kai paprastas ir refleksyvus pasitikėjimas susilieja į vieningą visumą.

Kiti mokslininkai pabrėžia neigiamus DI poveikio aspektus. Y. Chen et al. (2022) tyrime atskleidžiama, kad DI integracija į dalijimosi ekonomikos platformas, pavyzdžiui, „Airbnb“, gali sumažinti klientų įsitraukimą ir lojalumą dėl žmogiškosios sąveikos trūkumo. Šie rezultatai pabrėžia, kad neapčiuopiamas DI pobūdis ir vartotojų pirmenybė tiesioginiam kontaktui gali trukdyti pasitikėjimo kūrimui bei vystymui. Be to, Larsson'o ir Heintz'o (2020) tyrime išskiriamas skaidrumo vaidmuo kaip esminis pasitikėjimo veiksnys, siūlant įtraukti teisinių, socialinių ir humanitarinių sričių žinias į DI valdymo procesus.

Galiausiai, Ulfert et al. (2024) bei Yang ir Rau'as (2024) nagrinėja pasitikėjimo dinamiką žmogaus ir DI sąveikoje. Ulfert et al. (2024) išskiria, kad pasitikėjimas DI komandose yra dvikryptis, apimantis ne tik žmogaus, bet ir DI gebėjimą vertinti kitų narių patikimumą. Tuo tarpu Yang ir Rau'as (2024) nustatė, jog pradinis pasitikėjimas DI labiau veikia galutinį pasitikėjimą nei pradinis pasitikėjimas žmogumi, tačiau žmogaus teikiamas grįžtamasis ryšys yra svarbesnis pasitikėjimo vystymui ir emociniam sužaditimui. Šie tyrimai atskleidžia, kad pasitikėjimas DI yra sudėtingas ir daugialypis reiškinys, priklausantis nuo socialinių, techninių ir emocinių veiksnių.

Pasitikėjimo žmogaus kuriamu turiniu tematika moksliniuose darbuose yra rečiau pasitaikanti nei tyrimai, sutelkti į pasitikėjimą DI, vertinant straipsnius, išleistus per pastaruosius penkis metus. Vis dėlto, atsižvelgiant į skirtingas mokslininkų naudojamas metodologijas, analizuojamus kontekstus, tyrimai, sutelkti į žmogaus generuojamą turinį, prieina skirtingus rezultatus – panašus pastebėjimas buvo padarytas vertinant straipsnius, orientuotus į pasitikėjimą dirbtiniu intelektu bei jo kuriamu turiniu. Toliau pateikiamas keleto tyrimų apie pasitikėjimą žmogaus kuriamu turiniu apibendrinimas, kuris atskleidžia įvairius mokslinius požiūrius ir išvadas.

Zhang et al. (2023) tyrimas nagrinėja pasitikėjimo socialiniais tinklais apibrėžimus, veiksnus ir pasekmes, remiantis sistemine 70 mokslinių darbų apžvalga. Tyrime analizuota, kaip pasitikėjimas socialiniais tinklais buvo konceptualizuotas, matuojamas ir koks jo poveikis vartotojų elgesiui bei požiūriui. Rezultatai parodė, kad tik apie pusė tyrimų pateikė aiškius pasitikėjimo apibrėžimus, daugiausia susijusius su vartotojų pasirengimu pilnai pasikliauti kitų socialiniuose tinkluose dalyvaujančių vartotojų skelbiamu turiniu, neieškant šaltinių ir kitų būdų patikrinti informacijos validumą. Pastebėta, jog pasitikėjimą lemia demografiniai veiksniai, pavyzdžiui, amžius, rasė, išsilavinimas; sąveika su socialiniais tinklais, pavyzdžiui, naudojimosi intensyvumas, susipažinimas su informacijos šaltiniu; ir pačių socialinių tinklų savybės, pavyzdžiui, skaidrumas, patikimumas. Pasitikėjimas socialiniais tinklais skatina vartotojus dalintis informacija, tačiau mažina poreikį tikrinti informacijos patikimumą.

Alkhamees'as et al. (2021) savo tyrime nagrinėjo naudotojų pasitikėjimą socialiniuose tinkluose, sistemingai apžvelgiant 199 straipsnius, iš kurių 70 buvo išsamiai išnagrinėti. Autoriai pastebėjo, jog moksliniuose darbuose pasitikėjimo tematika pagrindinis dėmesys skiriamas įvairioms temoms, tokioms kaip netikrų paskyrų atpažinimas, melagienių (angl. *fake news*) identifikavimas ir žmogaus sukurto turinio patikimumo vertinimas. Pasitikėjimą socialiniuose tinkluose lemia naudotojų kognityvinės, elgesio ir reputacinės savybės, taip pat turinio savybės ir struktūriniai profilio elementai. Tyrimo autoriai pabrėžia, kad nėra universalios priemonės, galinčios visapusiškai įvertinti vartotojų pasitikėjimą socialiniuose tinkluose. Dauguma sprendimų yra efektyvūs tik konkrečiomis

sąlygomis ir dar neapėrią įvairiapusio pasitikėjimo vertinimo, pritaikomo visiems socialiniams tinklams.

Hochstein et al. (2023) tyrimas nagrinėja vartotojų pasitikėjimą skaitmeninėje erdvėje ir jo vaidmenį žmogaus sukurtu turinio (angl. *user-generated content*) veiksmingumui. Autoriai teigia, kad vartotojų pasitikėjimas žmogaus kuriamu turiniu priklauso nuo dviejų pagrindinių sąlygų: turinio iškraipymo ir prarastos informacijos lygio. Paskyrų autentiškumo patvirtinimo funkcijos stiprina, tuo tarpu nuotraukų filtrai bei turinio pašalinimo funkcijos silpnina pasitikėjimą ir mažina žmogaus kuriamo turinio veiksmingumą. Taip pat atskleista, jog platformų moderavimo veikla gali sumažinti neribojamo turinio, pavyzdžiui, atsiliepimų, efektyvumą, tačiau padidina riboto turinio, pavyzdžiui, patiktukų, poveikį. Vartotojų pasitikėjimas skaitmeninėje erdvėje yra esminis norint užtikrinti, kad žmogaus kuriamas turinys teigiamai veiktų vartotojų pirkimo intencijas.

Mohamed et al. (2023) tyrimas analizuoja žmogaus sukurtu turinio (angl. *user-generated content*) patikimumo ir įsisavinimo įtaką vartotojų pirkimo ketinimams kosmetikos rinkoje Malaizijoje. Rezultatai atskleidė, kad informacijos patikimumas ir vartotojo gebėjimas priimti pateiktą informaciją stipriai veikia jų sprendimus dėl pirkimo. Pavyzdžiui, „YouTube“ platformoje pateikiami grožio nuomonės formuotojų atsiliepimai laikomi patikimais, didina vartotojų pasitikėjimą ir skatina juos priimti sprendimą pirkti vieną ar kitą rekomenduojamą produktą. Tyrimas parodė, kad informacijos įsisavinimas stiprina vartotojų gebėjimą pritaikyti gautas žinias praktikoje, o tai skatina vartotojų intenciją pirkti. Be to, naudotojų sąveika su gerą reputaciją turinčių nuomonės formuotojų turiniu formuoja didesnę įsitraukimą ir lojalumą produktams, todėl skaitmeninėse platformose patikimumas tampa esminiu vartotojų elgsenos formavimo elementu.

3 lentelė. Pasitikėjimas žmogaus sukurtu turiniu socialiniuose tinkluose

Autorius(-iai), metai	Tyrimo objektas	Tyrimo tikslai / hipotezės	Tyrimo rezultatai
Zhang et al. (2023)	Pasitikėjimas socialiniais tinklais, pasitikėjimo apibrėžimas, konceptualizavimas, matavimas, pasitikėjimo veiksniai bei jo įtaka vartotojų elgsenai bei ketinimams.	Pateikti sisteminę apžvalgą apie pasitikėjimą socialiniais tinklais, identifikuoti pagrindines temas, iššūkius ir pasiūlyti kryptis būsimiems tyrimams. Tyrimo klausimai: 1) kaip ankstesni tyrimai konceptualizavo ir apibrėžė pasitikėjimą socialiniais tinklais? 2) kaip empiriniu būdu buvo matuojamas pasitikėjimas? 3) kokie yra pasitikėjimo veiksniai socialiniuose tinkluose? 3) kokia yra pasitikėjimo įtaką vartotojų ketinimams socialiniuose tinkluose?	Tyrėjai, išanalizavę skirtingo pobūdžio tyrimus apie pasitikėjimą socialiniais tinklais, pastebėjimo, jog dauguma tyrimu naudojosi bendriniais klausimais, neapimančiais matavimų, atspindinčių specifinius aspektus. Pasitikėjimą lemia demografiniai aspektai, tokie kaip amžius, rasė, išsilavinimas, taip pat sąveika su socialiniais tinklais, pavyzdžiui, naudojimosi intensyvumas, pažintis su informacijos šaltiniu, bei pačių socialinių tinklų savybės, pavyzdžiui, skaidrumas ar socialinis patikimumas. Pasitikėjimas socialiniuose tinkluose talpinamu turiniu skatina vartotojus dalintis informacija, naudotis socialiniais tinklais, tačiau sumažina vartotojo suvokiamą poreikį tikrinti informacijos patikimumą.
Alkhamees et al. (2021)	Vartotojų pasitikėjimo	Tyrimo tikslai:	Vartotojų pasitikėjimas socialiniuose tinkluose gali būti

Autorius(-iai), metai	Tyrimo objektas	Tyrimo tikslai / hipotezės	Tyrimo rezultatai
	matavimas socialiniuose tinkluose	1) sistemingai peržiūrėti tyrimus apie naudotojų pasitikėjimą socialiniuose tinkluose. 2) identifikuoti pagrindinius pasitikėjimo veiksnius ir charakteristikas. 3) išnagrinėti aplinkybes, kuriose buvo analizuojamas pasitikėjimas. 4) išskirti tyrimų spragas ir pasiūlyti ateities kryptis.	apibrėžiamas per kelias pagrindines charakteristikas: pažintinės charakteristikos apima naudotojų žinias ir gebėjimus specifinėse srityse, tuo tarpu elgesio charakteristikos apima naudotojų aktyvumą ir sąveiką su kitais tinklo dalyviais. Straipsnio autoriai išskiria, kad tyrimai pasitikėjimo socialiniuose tinkluose tematika, atliekami įvairiomis aplinkybėmis, analizuojant netikras socialinių tinklų paskyras, melagingus pranešimus, rekomendacijas. Visgi, išskiriama, jog daugeliu atveju tyrimai susiduria su mažų ar ribotų duomenų rinkinių iššūkiais, naudoja tik vieną algoritmą. Autorių teigimu, tolimesni tyrimai geresnių rezultatų galėtų pasiekti į savo metodologiją įtraukdami giluminius mokymosi algoritmus, kurie leistų efektyviai apdoroti didelius duomenų rinkinius. Reikšmingai pagelbėtų ir platesnis metodų bei duomenų taikymas.
Hochstein et al. (2023)	Vartotojų pasitikėjimas skaitmeninėje erdvėje bei jo vaidmuo vertinant žmogaus sukurtą turinio (angl. <i>user-generated content</i>) veiksmingumą bei įtaką įmonių rezultatams.	Tyrimas siekia atliepti tris tikslus: 1) sukurti teoriją apie vartotojų pasitikėjimą skaitmeninėje erdvėje, išryškinant sąlygas, būtiną, kad vartotojai pasitikėtų žmogaus sukurtu turiniu; 2) pateikti įrodymus, kaip tam tikros platformų funkcijos gali padidinti arba sumažinti vartotojų pasitikėjimą. 3) išanalizuoti, kaip pasitikėjimo skaitmeninėje erdvėje sąlygos veikia vartotojų elgesį ir įmonių rezultatus. Tyrimo hipotezės: 1) naudotojų sukurtas turinys teigiamai veikia įmonių rezultatus. 2) funkcijos, kurios užtikrina paskyros autentiškumą (pavyzdžiui, patvirtinimai), stiprina patikimumą bei teigiamą turinio poveikį. 3) funkcijos, kurios leidžia keisti turinį (pavyzdžiui, foto filtrai), silpnina vartotojų pasitikėjimą bei teigiamą turinio poveikį. 4) funkcijos, kurios leidžia ištrinti turinį, mažina vartotojų pasitikėjimą turiniu ir teigiamą jo poveikį.	Siekiant didinti vartotojų pasitikėjimą, būtina užtikrinti, kad platformoje talpinamas turinys nėra klaidinantis ar nepilnas. Šios sąlygos svarbios siekiant ugdyti vartotojų pasitikėjimą bei skatinant juos priimti sprendimus, kurie būtų grindžiami turiniu. Skaitmeninėje erdvėje esančių platformų funkcionalumas daro itin didelę įtaką vartotojų pasitikėjimui. Funkcijos, kurios užtikrina turinio kūrėjų paskyrų autentiškumą, pavyzdžiui, patvirtinimai ar „patikimumo ženklai“, skatina vartotojų pasitikėjimą ir stiprina teigiamą turinio poveikį įmonių veiklai. Tuo tarpu funkcijos, leidžiančios keisti turinį, tokios kaip foto filtrai, mažina pasitikėjimą, nes gali sukelti įtarimų dėl pateikiamos informacijos tikrumo. Be to, turinio ištrynimo galimybė dažnai sukelia įtarimus dėl prarastos informacijos, o tai trukdo vartotojams sudaryti informuotas išvadas, vertinant turinį. Platformos moderavimo aprėptis žiniasklaidoje taip pat veikia pasitikėjimą: intensyvus dėmesys platformų moderavimo

Autorius(-iai), metai	Tyrimo objektas	Tyrimo tikslai / hipotezės	Tyrimo rezultatai
		5) platformos moderavimo žiniasklaidos aprėptis gali silpninti naudotojų pasitikėjimą turiniu.	veiklai mažina mažai ribojamo turinio, tokio kaip komentarai ar įrašai, patikimumą, tačiau gali sustiprinti kitokio pobūdžio turinio, tokio kaip patiktukai, poveikį. Tyrimas pabrėžia, kad neatkreipiant dėmesio į šias skaitmeninio pasitikėjimo sąlygas, turinys praranda savo efektyvumą tiek vartotojų elgesio formavime, tiek įmonių rezultatų gerinime.
Mohamed et al. (2023)	Žmogaus kuriamo turinio patikimumo ir įsisavinimo įtaka vartotojų ketinimams pirkti.	Išsiaiškinti, kaip žmogaus sukurtos informacijos patikimumas ir įsisavinimas veikia vartotojų pirkimo intencijas. Hipotezės: 1) informacijos patikimumas teigiamai veikia pirkimo ketinimus 2) informacijos įsisavinimas teigiamai veikia pirkimo ketinimus	Tiek informacijos patikimumas, tiek jos įsisavinimas skatina vartotojus pirkti. Informacijos patikimumas skatina vartotojų pasitikėjimą ir didina jų polinkį priimti sprendimą pirkti, tuo tarpu informacijos įsisavinimas leidžia vartotojams ne tik įsisavinti informaciją, bet ją pritaikyti praktikoje, suvokiant kaip vienas ar kitas produktas galėtų spręsti vieną ar kitą problemą. Nustatyta, kad nuomonės formuotojai, dalindamiesi savo pastebėjimais apie produktus, skatina turinio patikimumą, taip skatinant vartotojus priimti sprendimą pirkti. Žmogaus sukurtos informacijos kokybė ir jos įsisavinimas yra esminiai veiksniai, skatinantys vartotojų pirkimo ketinimus. Socialinių tinklų platformos, tokios kaip „YouTube“, yra efektyvus kanalas, leidžiantis nuomonės formuotojams skleisti patikimą informaciją, kuri formuoja vartotojų elgesį.

Remiantis įvairių mokslininkų darbų analize, pasitikėjimas žmogaus kuriamu turiniu yra sudėtingas reiškinys, priklausantis nuo konteksto ir analizuojamų veiksnių. Mohamed et al. (2023) nustatė, kad patikimumas ir informacijos įsisavinimas skatina vartotojų pirkimo sprendimus, o Hochstein et al. (2023) parodė, kad paskyrų patvirtinimai didina pasitikėjimą, tuo tarpu nuotraukų filtrai ir turinio pašalinimas jį mažina. Zhang et al. (2023) pabrėžė pasitikėjimo veiksnių įvairovę, tokių kaip demografiniai rodikliai ir tinklų skaidrumas, tačiau atkreipė dėmesį į tyrimų apibrėžimų ir matavimų ribotumą. Alkhamees et al. (2021) pastebėjo, kad nėra universalių priemonių, galinčių visapusiškai įvertinti žmogaus kuriamo turinio patikimumą. Tyrimai parodo, kad pasitikėjimas žmogaus kuriamu turiniu yra daugialypis procesas, kurį būtina nagrinėti toliau, siekiant geriau suprasti vartotojų elgseną ir kurti veiksmingas pasitikėjimo stiprinimo strategijas.

Nors akademinio darbo, orientuoto į pasitikėjimą DI ir žmogaus kuriamu turiniu atskirai yra ganėtinai nemažai, pastebimas tyrimų, kurie lygintų šių dviejų skirtingų autorių turinį, trūkumas. Tuo tarpu

atkreipiant dėmesį į tai, jog tobulėjant mašininio mokymosi modeliams, ypač didžiosios kalbos modeliams (angl. *Large Language Models* arba *LLMs*), tokiems kaip GPT, atskirti žmogaus ir DI sukurtą turinį vieną nuo kito darosi vis sudėtingiau. Xie et al. (2024) teigia, jog tradiciniai DI aptikimo įrankiai, tokie kaip loginė regresija, efektyviai atskiria DI ir žmogaus sukurtą tekstinį turinį. Įžvelgiami dideli skirtumai tarp žmogaus ir DI kuriamo turinio per naudojamą stilių: DI labiau linkęs kurti su žala susijusį turinį, tuo tarpu žmogaus kūryba dažniau sutelkiama į moralines vertybes. Anot Gu et al. (2024) didesnis dėmesys realistiškumui ir fantazijai kuriant vaizdinį turinį DI įrankių pagalba suteikia reklamoms natūralesnį ir kūrybiškesnį įvaizdį. DI kuriamas turinys gali sukelti nejaukumą, kai atvaizduojami vaizdiniai elementai atrodo nenatūraliai, yra ne vietoje ar nėra sulieti su juos supančiu kontekstu. Lyginamosios analizės tarp dirbtinio intelekto ir žmogaus kuriamo turinio neretais atvejais analizuoja išskirtinai tekstinę informaciją. Vis dėlto, dirbtinis intelektas yra itin pažengęs ne tik teksto, bet ir vaizdo, garso, meno kūrimo srityse (Park et al., 2024).

Toliau pateikiamoje lentelėje (žr. 4 lentelę) išskiriamos pagrindinės tematikos, aktualios magistro baigiamajam projektui bei susijusios su dirbtinio intelekto ir žmogaus kuriamu turiniu.

4 lentelė. Pagrindinės tematikos moksliniuose darbuose apie pasitikėjimą DI ir žmogaus kuriamu turiniu

Tematika	Subtemos	Pagrindinės įžvalgos	Literatūros pavyzdžiai
Turinio tipai ir savybės	- Kūrėjas (žmogus, DI); - turinio forma (tekstas, vaizdas, garsas, vaizdo įrašai); - autentiškumas; - skaidrumas; - autorystė; - šaltinio patikimumas.	DI generuojamas turinys yra sunkiai atskiriamas nuo žmogaus kuriamo. DI labiau linkęs į struktūruotą ir realistišką turinį, tačiau jam sunku būti autentiškam. Žmogaus turinys yra pasižymi kontekstualumu ir morališkumu. Šiame kontekste patikimumas yra tiesiogiai susijęs su reputacija ir emociniu ryšiu. Skaidrumas ir autorystė svarbūs abiem turinio tipams.	(Gu et al., 2024; Larsson ir Heintz, 2020; Mohamed et al., 2024; Park et al., 2024; Xie et al., 2024)
Vartotojų požiūris	- Pasitikėjimas DI kuriamu turiniu; - pasitikėjimas žmogaus kuriamu turiniu; - vartotojų pasitikėjimas turiniu.	Pasitikėjimas DI priklauso nuo pradinio įvertinimo ir skaidrumo. Žmogaus kuriamam turiniui didelę įtaką daro autentiškumas, reputacija ir emocinis ryšys.	(Mohamed et al., 2024; Scharowski et al., 2023; Zhang et al., 2023)
Pasitikėjimą lemiantys veiksniai	- Episteminiai ir pragmatiniai aspektai; - žmogiškoji sąveika; - techninių sprendimų paaiškinimai.	DI gebėjimas aiškinti sprendimus padidina pasiklovimą. Skaidrumo lygis tiesiogiai veikia pasitikėjimą.	(Ferrario et al., 2020; Ulfert et al., 2024)
Pasitikėjimo poveikis ketinimams	- Vartotojų elgsena; - informacijos dalinimasis.	DI gali skatinti pasiklovimą, tačiau žmogiškoji sąveika dažniau veikia emocinį ryšį ir pasitikėjimą.	(Y. Chen et al., 2022; Hochstein et al., 2023; Yang ir Rau, 2024)

Dirbtinio intelekto ir žmogaus kuriamas turinys pasižymi skirtingais tipais ir savybėmis, kurios daro įtaką vartotojų pasitikėjimui. Keletas iš dažniausiai naudojamų turinio tipų apima tekstą, vaizdus, garsus ir vaizdo įrašus. DI generuojamas turinys yra struktūruotas, realistiškas ir greitai kuriamas, tačiau jam dažnai trūksta autentiškumo bei skaidrumo (Larsson ir Heintz, 2020; Park et al., 2024). Tuo tarpu žmogaus kuriamas turinys išsiskiria stipresniu emociniu patrauklumu ir moralinėmis vertybėmis, kurios padidina pasitikėjimą. Šio turinio patikimumą nulemia autorystė ir reputacija

(Mohamed et al., 2024; Zhang et al., 2023). Tobulėjant DI modeliams, turinio tipai darosi vis panašesni, o tai kelia iššūkių užtikrinant skaidrumą ir autorystės atpažinimą (Xie et al., 2024).

Vartotojų požiūris į dirbtinio intelekto ir žmogaus kuriamą turinį yra formuojamas įvairių socialinių, techninių ir emocinių aspektų. DI sukurtas turinys sulaukia daugiau pradinio pasitikėjimo dėl savo tikslumo ir efektyvumo, tačiau pasikliovimas juo priklauso nuo paaiškinimų apie sprendimų priėmimo procesą ir skaidrumą (Scharowski et al., 2023). Pasitikėjimas yra lemiamas vartotojų požiūrio elementas, kuris veikia jų ketinimus naudotis turiniu, priimti rekomendacijas ar dalintis informacija socialiniuose tinkluose. Pavyzdžiui, DI gebėjimas paaiškinti sprendimus gali padidinti pasikliovimą tik esant aiškioms ir suprantamoms rekomendacijoms, o žmogaus kuriamas turinys labiau paveikia vartotojus per emocinį sužadimą ir interaktyvumą (Yang ir Rau, 2024). Tokiu būdu vartotojų požiūris tampa svarbiu faktoriumi analizuojant abiejų turinio tipų poveikį vartotojų pasitikėjimui ir ketinimams socialiniuose tinkluose.

Pasitikėjimo dirbtiniu intelektu ir žmogaus kuriamu turiniu formavimasis yra sudėtingas procesas, kuriame svarbų vaidmenį atlieka skaidrumas, patikimumas ir socialinis kontekstas. Ferrario et al. (2020) siūlo „inkrementinį pasitikėjimo modelį“, kuris pabrėžia paprasto, refleksyvaus ir paradigminio pasitikėjimo sluoksnius, atskleidžiančius, kaip episteminiai (žiniomis pagrįsti) ir pragmatiniai (naudos siekimo) veiksniai prisideda prie vartotojų pasitikėjimo. Ulfert et al. (2024) nustatė, kad pasitikėjimą DI komandose lemia ne tik technologinis skaidrumas, bet ir DI gebėjimas pagrįsti savo sprendimus, parodyti integralumą bei geranoriškumą. Be to, skaidrumo vaidmuo pabrėžia multidisciplininio požiūrio svarbą, nes įtraukiant socialines ir teises ekspertizes galima ne tik užtikrinti patikimumą, bet ir sumažinti galimas etines ar socialines rizikas. Tokiu būdu pasitikėjimo veiksmų analizė tampa būtina stiprinant DI ir žmogaus kuriamo turinio sąveiką.

Įvairūs pasitikėjimą formuojantys aspektai gali labai paveikti vartotojų elgseną bei sprendimus. Hochstein et al. (2023) teigia, jog pasitikėjimas yra pagrindinis elementas, lemiantis vartotojų pirkimo sprendimus. Pasitikėjimą didina papildomos skaitmeninės funkcijos, tokios kaip autentiškumo žymėjimai ir patikimumo ženklai. Yang ir Rau'as (2024) pabrėžia emocinę sąveikos svarbą – nors DI pradinė įtaka pasitikėjimui yra reikšminga, žmogaus teikiamas grįžtamasis ryšys labiau veikia pasitikėjimo formavimą ir sužadina teigiamas emocijas, stiprinančias vartotojų ryšį su turiniu. Tuo tarpu Y. Chen et al. (2022) pažymi, kad DI integracija gali veikti kaip neigiamas moderatorius tarp pasitikėjimo platformomis ar šeiminiškais ir vartotojų lojalumo bei įsitraukimo.

Peržvelgus akademinius darbus, pastebima, kad tyrimų apie vartotojų požiūrį į dirbtinio intelekto ir žmogaus kuriamą vizualinį ir tekstinį turinį socialiniuose tinkluose vis dar trūksta, ypač nagrinėjant šio turinio poveikį vartotojų pasitikėjimui ir ketinimams. Nors yra mokslinių darbų, nagrinėjančių pasitikėjimą dirbtinio intelekto ir žmogaus kuriamu turiniu, vis dar trūksta išsamių tyrimų apie tai, kokią įtaką įvairios turinio savybės, tokios kaip autentiškumas, autorystė ar šaltinio patikimumas, daro vartotojų elgesiui ir sprendimams. Pasitikėjimas yra esminis sėkmingos sąveikos tarp žmogaus ir DI elementas. Praktiškai tai reiškia, kad tyrimai apie pasitikėjimą DI kontekste gali padėti gerinti komandų efektyvumą, sumažinti klaidas ir skatinti DI pritaikymą įvairiose organizacijose. Be to, reikėtų atsižvelgti į tai, kaip šios savybės veikia emocinį vartotojų įsitraukimą ir ilgalaikį lojalumą platformoms ar prekių ženkams. Ši spraga pabrėžia poreikį gilintis į šiuos aspektus, siekiant pateikti nuodugnesnę vartotojų elgsenos analizę, vertinant tiek tekstinį, bet ir vizualinį turinį. Papildomai svarbu tirti dirbtinio intelekto generuojamo turinio skaidrumą ir paaiškinamumą, kurie daro įtaką vartotojų pasitikėjimui ir norui naudotis tokiomis technologijomis. Tokie tyrimai praturtintų mokslinę

literatūrą, suteiktą naudingų įžvalgų turinio kūrėjams ir socialinių tinklų platformų administratoriams apie tai, kaip stiprinti vartotojų pasitikėjimą, skatinti teigiamus sprendimus ir sukurti etišką bei į žmogų orientuotą socialinių tinklų ekosistemą.

Apibendrinant visus šioje magistro baigiamojo darbo dalyje analizuotus mokslinius darbus, jose pateikiamus rezultatus bei rekomendacijas tolimesniems tyrimams, išskiriamos probleminės sritys, kurias reikia tirti išsamiau, ieškant atsakymo į žemiau įvardintus klausimus:

- kaip dirbtinio intelekto ir žmogaus kuriamas vizualinis bei tekstinis turinys socialiniuose tinkluose veikia vartotojų pasitikėjimą bei ketinimus?;
- kaip vartotojo pasitikėjimą formuoja dirbtinio intelekto ir žmogaus kuriamo turinio savybės, tokios kaip autentiškumas, autorystė ir kokybė?

2. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams tyrimo teorinis pagrindimas

Šioje magistro projekto dalyje gilinamasi į dirbtinio intelekto ir žmogaus kuriamo turinio poveikio vartotojų pasitikėjimui ir ketinimams teorinį pagrindimą. Aptariami dirbtinio intelekto ir žmogaus sukurtų sprendimų skirtumai, autorystės atskleidimo poveikis bei turinio autentiškumo ir skaidrumo svarba formuojant vartotojų pasitikėjimą. Skyrius apibendrinamas mokslinės literatūros analize grindžiamu konceptualiu modeliu, kuris atskleidžia sąsajas tarp DI ir žmogaus kuriamo turinio ypatybių bei jų įtakos vartotojų ketinimams socialiniuose tinkluose.

2.1. DI ir žmogaus kuriamo turinio socialiniuose tinkluose charakteristikos

Socialiniai tinklai iš esmės transformavo skaitmeninę komunikaciją. Šiame skyrelyje nagrinėjama DI ir žmogaus kuriamo turinio savybės bei jų vaidmuo socialinių tinklų komunikacijoje.

2.1.1. Socialiniuose tinkluose kuriamo komunikacinio turinio tipai bei savybės

Socialiniai tinklai tapo viena svarbiausių skaitmeninės komunikacijos priemonių, keičiančių, kaip žmonės bendrauja tarpusavyje ir kaip verslai pasiekia savo klientus. Socialinių tinklų platformos suteikia galimybę kurti įvairialypį turinį, kuris gali būti pritaikytas siekiant patenkinti įvairių tikslinių auditorijų poreikius. Socialiniai tinklai įgalina vartotojus tapti aktyviais dalyviais per vartojimą, įsitraukimą ir kūrybą, susijusią su prekių ženklais. Turinio vartojimas socialiniuose tinkluose, pavyzdžiui, reklamų peržiūra ar produktų apžvalgų skaitymas, padeda didinti prekės ženklo žinomumą ir įtraukti naujus vartotojus. Prisidėjimas, pavyzdžiui, komentarų rašymas, įvertinimų teikimas, skatina organišką turinio plitimą, o kūryba, pavyzdžiui, naudotojų kuriamo turinio paskelbimas, leidžia vartotojams aktyviai dalyvauti prekės ženklo komunikacijoje ir formuoti jo įvaizdį (Buzeta et al., 2020). Be to, socialiniai tinklai pasižymi savybėmis, kurios didina jų efektyvumą: interaktyvumas – vartotojai dalijasi savo nuomonėmis, reaguoja į turinį ir aktyviai jį kuria; personalizacija – turinys gali būti pritaikytas specifiniams vartotojų poreikiams; greitis – informacija skleidžiama ir vartojama itin sparčiai, o įvairiapusiškumas leidžia naudoti tekstus, vaizdus, vaizdo įrašus ir kitus formatus (Schreiner et al., 2021). Tokiu būdu efektyvus socialinių tinklų panaudojimas gali ne tik padidinti prekių ženklo matomumą, bet ir kurti ilgalaikį vartotojų lojalumą.

Shahbaznezhad'as et al. (2021) išskyrė tris turinio socialiniuose tinkluose kategorijas: racionalus (angl. *rational*, *informational*), emocionalus (angl. *emotional*) bei transakcinis (angl. *transactional*). Autoriai racionalų turinį apibūdino kaip informuojantį, edukuojantį, sprendžiantį vartotojų problemas, emocinis turinys siekia sužadinti emocijas, įtraukti per pramogas ar vizualų patrauklumą, o transakcinis turinys, anot autorių, skirtas skatinti tiesioginius veiksmus, pavyzdžiui, pirkimą ar paslaugų užsakymą.

Remiantis „Statistia“ duomenimis, populiariausi socialiniai tinklai Lietuvos gyventojų tarpe 2023 metais buvo „Facebook“, „Instagram“, „Pinterest“, „YouTube“, „X“. Toliau pateikiama apibendrinamoji lentelė (žr. 5 lentelę), apimanti Lietuvos gyventojų tarpe populiariausius socialinius tinklus bei juose publikuojamo turinio pobūdį.

5 lentelė. Populiariausi socialiniai tinklai Lietuvoje bei jų turinio formatų charakteristikos

Socialinis tinklas	Pagrindiniai turinio formatai	Dažniausiai dominuojančios kategorijos	Charakteristikos
„Facebook“	Tekstai, nuotraukos, vaizdo įrašai, apklausos	Racionalus, emocinis, transakcinis	Universali platforma, tinkama įvairaus tipo turiniui; aukštas interaktyvumo lygis.
„Instagram“	Nuotraukos, trumpi vaizdo įrašai („Reels“), istorijos („Stories“)	Emocinis, transakcinis	Orientuota į estetinį, vizualiai patrauklų turinį; didelis jaunimo aktyvumas.
„Pinterest“	Nuotraukos, infografika, vizualinės kolekcijos	Racionalus, emocinis	Populiari idėjų ir įkvėpimo platforma.
„YouTube“	Ilgi vaizdo įrašai, trumpi vaizdo klipai („YouTube Shorts“), tiesioginės transliacijos	Racionalus, emocinis	Populiariausia vaizdo įrašų platforma; puikiai tinka edukaciniam ir pramoginiam turiniui.
„X“	Trumpi tekstai, vaizdo įrašai, nuotraukos	Racionalus, transakcinis	Efektyvus greitai informacijos sklaidai ir aktualijoms; dažnai naudojamas diskusijoms.

„Facebook“, Lietuvos gyventojų tarpe, yra populiariausia socialinių tinklų platforma, pasižyminti savo universalumu, tinkanti tiek informaciniam, tiek emociniam bei transakciniam turiniui. „Instagram“ orientuojasi į vizualų turinį, o jo vartotojai labiau linkę reaguoti į estetiškai patrauklius įrašus. „YouTube“, kaip ilgųjų vaizdo įrašų platforma, leidžia gilinti auditorijos įsitraukimą per edukacinį ar pramoginį turinį. „X“ skatina trumpas ir aiškias žinutes, o „Pinterest“ – įkvėpimo ir idėjų dalijimąsi per vizualias kolekcijas. Socialinių tinklų populiarumas ir skirtingų platformų specifika leidžia kiekvienai iš jų tapti efektyvia komunikacijos priemone, pritaikyta įvairiems tikslams ir auditorijoms. Kadangi skirtingi tinklai turi unikalius funkcionalumus, turinio kūrėjams svarbu suprasti, kaip geriausiai išnaudoti kiekvienos platformos galimybes, norint pasiekti didžiausią vartotojų įsitraukimą ir įgyvendinti komunikacinius tinklus.

Emocijų poveikis yra esminis socialinių tinklų strategijų elementas. Kaip pastebėjo Schreiner et al. (2021), turinys, kuris sukelia teigiamas emocijas, padidina vartotojų įsitraukimą, ypač kai naudojamas turtingas medijos formatas (angl. *Media richness*), tokie kaip vaizdo įrašai. Emocinis poveikis ypač stiprus platformose, kur auditorijos aktyvumas yra aukštas. Pavyzdžiui, „Instagram“ ir „YouTube“ skatina vartotojus ne tik žiūrėti, bet ir reaguoti per komentarus ar dalijimąsi. Medijos turtingumo teorija (angl. *Media Richness theory*) teigia, kad vaizdo įrašais ar animacijomis praturtintas turinys yra efektyvesnis emociniam poveikiui sukurti nei tekstinis (Shahbaznezhad et al., 2021). Tačiau emocionalaus turinio sėkmė labai priklauso nuo platformos specifikos ir auditorijos poreikių. Platformose, tokiose kaip „Pinterest“, emocinis turinys dažnai siejamas su estetiniu pasitenkinimu, tuo tarpu „YouTube Shorts“ didesnę dėmesį skiria greitai sužadinančioms emocijoms.

Remiantis Dabbous ir Barakat (2020) bei Yeung'o et al. (2022) tyrimais, turinio kokybiškumą galima įvertinti atsižvelgiant į kelis pagrindinius kriterijus: tikslumą ir patikimumą, gebėjimą įtraukti, personalizotumą, vizualinį patrauklumą ir suprantamumą, aiškų nurodymą, kokių veiksmų iš vartotojo tikimasi bei inovatyvumą. Toliau pateikiamoje lentelėje (žr. 6 lentelę) atsispindi pagrindiniai aukščiau minėti kokybiško turinio kriterijai bei juos apibūdinančios charakteristikos.

6 lentelė. Turino kokybės vertinimas

Kriterijus	Apibūdinančios charakteristikos
Tikslumas	Faktų pagrįstumas, klaidingos informacijos vengimas.
Gebėjimas įtraukti	Auditorijos reakcijų (komentarų, dalinimosi) skatinimas.
Personalizacija (aktualumas)	Turinio pritaikymas specifinei auditorijai.
Vizualinis patrauklumas ir suprantamumas	Estetiškai maloni, lengvai suprantama turinio pateiktis.
Veiksmo iš vartotojo skatinimas	Informacijos, kaip elgtis toliau, pateikimas, pavyzdžiui, kreiptis į specialistus, registruotis į koncertą.
Inovatyvumas	Kūrybiniai sprendimai, išskiriantys turinį iš kitų.

Tikslumas yra kertinis kokybiško turinio požymis, užtikrinantis, kad informacija yra pagrįsta faktais, paruošta, remiantis patikimais šaltiniais. Klaidinga informacija socialiniuose tinkluose gali būti labai paplitusi ir paveiki. Tikslus turinys kuria pasitikėjimą ir leidžia prekės ženklui išsiskirti rinkoje kaip patikimam informacijos šaltiniui. Turinio **gebėjimas įtraukti** nurodo jo gebėjimą rezonuoti su vartotojais, skatinant juos veikti – dalintis, komentuoti, reaguoti. Įtraukiantis turinys stiprina emocinį ryšį tarp vartotojo ir prekės ženklo. Pavyzdžiui, interaktyvios apklausos ar kvietimai dalyvauti diskusijose skatina organišką vartotojų įsitraukimą. **Personalizuotumas** padeda turiniui tiesiogiai atliepti specifinius individualaus auditorijos nario poreikius, vertybes bei interesus. **Vizualinis patrauklumas ir suprantamumas** yra svarbūs veiksniai, padedantys greitai pritraukti vartotojų dėmesį ir padidinti informacijos suvokimą. Net klaidingas turinys sulaukia didelio dėmesio, jei jis yra vizualiai patrauklus ir lengvai suprantamas, todėl būtina užtikrinti, kad kokybiškas turinys būtų pateikiamas paprasta ir vizualiai įtraukiančia forma. Turinys turėtų pateikti vartotojui aiškią informaciją apie tai, kokių **veiksmų tikimasi**, pavyzdžiui, užsiregistruoti, įsigyti produktą ar dalyvauti renginyje. Ši turinio savybė suteikia vartotojui aiškias instrukcijas, kokio veiksmo iš jo tikimasi. **Inovatyvumas** suteikia turiniui unikalumo, leidžiančio išsiskirti konkurencingame socialinių tinklų sraute. Kūrybiškai pateiktas turinys ne tik pritraukia dėmesį, bet ir kuria nepamirštamą vartotojo patirtį (Dabbous ir Barakat, 2020; Yeung et al., 2022).

2.1.2. DI panaudojimas socialiniuose tinkluose bei jo pagalba sukurto turinio ypatumai

Dirbtinio intelekto samprata ir taikymas socialiniuose tinkluose tapo viena iš aktualiausių technologinių temų pastaruosiu metu. DI apibrėžiamas kaip technologija, gebanti imituoti žmogaus intelektą ir atlikti kompleksines užduotis, apimančias mokymąsi, problemų sprendimą, kalbos atpažinimą ir kt. (Du-Harpur et al., 2020). Šios sistemos veikia naudojamos algoritmus ir didelius duomenų kiekius, kurie leidžia joms mokytis ir tobulėti savarankiškai. DI svarba šiuolaikinėje visuomenėje nuolat didėja dėl jo gebėjimo didinti žmogiškojo darbo našumą už itin mažus naudojimo kaštus (Kelly et al., 2023).

Prognozuojama, kad DI ir juo grįsti sprendimai pasaulio ekonomikai iki 2030-ųjų metų gali sugeneruoti papildomus 15 trilijonų JAV dolerių, tokiose srityse kaip sveikatos priežiūra, švietimas, transportas, klientų aptarnavimas ir rinkodara (Kelly et al., 2023). Pavyzdžiui, sveikatos srityje DI naudojamas ankstyvai ligų diagnostikai, švietime – individualizuotoms mokymo priemonėms kurti, transporto sektoriuje – automatizuotoms transporto priemonėms vystyti, o rinkodaroje – personalizavimui (Du-Harpur et al., 2020; Kelly et al., 2023).

Socialiniuose tinkluose dirbtinio intelekto taikymas tapo itin aktualus dėl potencialaus vartotojų patirties gerinimo ir įsitraukimo didinimo. Socialinių tinklų platformos, tokios kaip „Facebook“ ir „Instagram“, naudoja DI sistemas turinio personalizavimui, vartotojų elgsenos bei interesų analizei, siekiant pateikti vartotojui aktualiausią ir labiausiai tinkantį turinį (Novelli et al., 2024). Nors turinio personalizavimas pagal vartotojų paieškos istoriją, interneto naršymo pobūdį, ryšius socialiniuose tinkluose bei geografinius duomenis padidina vartotojų įsitraukimą, kyla iššūkių dėl informacijos įvairovės ir poliarizacijos. Algoritmai dažnai skatina „informacinio šulinio“, dar kitaip žinomo kaip „filtrų burbulo“ efektą (angl. *filter bubble*) – reiškinį apimantį situacijas, kai vartotojai mato tik jų pažiūras patvirtinantį turinį, taip apribojant galimybę susidurti su kitokiomis nuomonėmis ar nauja informacija. Pavyzdžiui, sveikatos kontekste tai gali nulemti, kad tėvai, skeptiškai žiūrintys į vakcinas, matys tik su jų įsitikinimais suderintą informaciją, t. y. vakcinų naudą neigiantį turinį, kas dar labiau sustiprins jų nusistatymą prieš skiepus (Holone, 2016).

Dirbtinis intelektas taip pat naudojamas turinio kūrimui automatizuoti, generuojant tekstus, vaizdus ir kitokio tipo skaitmeninį turinį (Du-Harpur et al., 2020). Šis procesas iš esmės pakeitė skaitmeninės kūrybos galimybes. Naudojant dirbtinio intelekto modelius, tokius kaip „ChatGPT“, socialinių tinklų vartotojai gali per trumpą laiką ir minimalius kaštus sugeneruoti tekstus, kurie atrodo natūraliai ir atitinka specifinius vartotojo poreikius. Pavyzdžiui, automatiškai kurti suasmenintus įrašus, optimizuotus reklamos tekstus ar net straipsnius pagal specifines temas, labiausiai atitinkančias vieno ar kito vartotojo poreikius. Dirbtinio intelekto pagalba automatiškai kuriamas ne tik tekstinis, bet ir vaizdinis turinys. Generatyviniai modeliai, tokie kaip „DALL-E“, automatizuotai kuria vizualinį turinį, kuris daugumai vartotojų atrodys estetiškai patrauklus ir profesionalus. Šie įrankiai turinio kūrėjams leidžia sumažinti kaštus, kurių reikėtų kuriant vaizdinį turinį tradiciniais būdais (Novelli et al., 2024). Nors tokie sprendimai gali pagerinti turinio kūrėjų produktyvumą, kyla klausimų dėl turinio autentiškumo ir originalumo. Dėl nuolatinio generatyvinio dirbtinio intelekto naudojimo kūryboje kyla rizika, kad turinys taps pasikartojantis ir praras unikalios žmogiškuosius niuansus, kurie tradiciškai laikomi kūrybos vertybe. Pavyzdžiui, automatiškai sugeneruotas vizualinis turinys gali neatitikti tikslinės auditorijos lūkesčių, jei ji kuriant nebuvo atsižvelgta į kultūrinį ar emocinį kontekstą. Be to, naudojant DI kūrybiniais tikslais, dažnai kyla etinių klausimų dėl autorinių teisių (Kelly et al., 2023; E. A. S. Mohamed et al., 2024b; Novelli et al., 2024).

Vartotojų pasitikėjimas DI generuojamu turiniu tampa vis svarbesniu aspektu socialiniuose tinkluose. Tyrimai rodo, kad vartotojams dažnai sunku atskirti, ar turinys sukurtas žmogaus, ar dirbtinio intelekto (Novelli et al., 2024; Park et al., 2024). Be to, DI sukurtas turinys gali būti panaudotas dezinformacijai skleisti ar kitais manipuliaciniais tikslais (Piduru, 2023a; Qi et al., 2024). Atliepiant tai, siekiama užtikrinti, kad dirbtiniu intelektu grindžiamų technologijų plėtra socialiniuose tinkluose būtų grindžiama aiškiomis etinėmis ir teisinėmis normomis (Kelly et al., 2023).

Turinio moderavimo srityje DI algoritmai padeda efektyviai aptikti netinkamą ar įžeidžiantį turinį, įskaitant dezinformaciją bei neapykantą kurstančius komentarus, o tai leidžia kurti saugesnę aplinką socialinių tinklų naudotojams. Visgi, svarbu atkreipti dėmesį, kad DI turinio moderavimo efektyvumas tiesiogiai priklauso nuo apmokomųjų duomenų kokybės ir algoritmų skaidrumo (Kelly et al., 2023). Automatizuotas moderavimas kartais susiduria su neteisingu konteksto interpretavimu, ar kitomis problemomis.

Socialiniuose tinkluose DI taip pat naudojamas vartotojų sąveikai gerinti, pavyzdžiui, per pokalbių robotus, padedančius atsakyti į vartotojų klausimus ir spręsti problemas. Šie robotai dažnai diegiami

į verslus, su tikslu gerinti vartotojo patirti per kokybišką klientų aptarnavimą (Du-Harpur et al., 2020). Nors jie didina aptarnavimo efektyvumą, kai kurie vartotojai mano, kad tokie robotai yra mažiau empatiški ir supratingi nei žmonės, negeba prisitaikyti prie kompleksinių situacijų.

Toliau pateikiamoje lentelėje (žr. 7 lentelę) išskiriami pagrindiniai dirbtinio intelekto funkcionalumai, taikomi socialiniuose tinkluose bei jų poveikis vartotojų patirčiai, kūrybai, požiūriui.

7 lentelė. Dirbtinio intelekto socialiniuose tinkluose sritys

Autorius(-iai), metai	DI sritis	Poveikis	Pastabos
Novelli et al. (2024), Piduru (2023)	Personalizavimas	Padidėjęs vartotojų įsitraukimas	Padidėjusi „informacinio šulinio“ rizika
Du-Harpur et al. (2020)	Automatinis turinio kūrimas	Efektyvesnė, greitesnė bei pigesnė turinio kūryba	Sumažėjęs kūrybinis originalumas
Kelly et al. (2023)	Turinio moderavimas	Nekokybiško turinio valdymas	Efektyvumas didėja, didėjant algoritmų skaidrumui
Holone (2016)	„Informacini šulinys“	Informacijos poliarizacija	Vartotojus pasiekiančio turinio įvairovės apribojimas
Mohamed et al. (2024)	Dirbtinio intelekto etika	Skatinamas etiškas DI naudojimas bei skaidrumas	Svarbi socialinė atsakomybė

Dirbtinio intelekto taikymas socialiniuose tinkluose iš esmės pakeitė kaip vartotojai sąveikauja su turiniu bei platformomis. Lentelėje atskleista tik maža dalis skaitmeninės erdvės sričių, kuriose naudojamas dirbtinis intelektas. DI grįsti įrankiai ypač populiarūs elektroninėse parduotuvėse ir paslaugų sektoriuje, kur naudojama papildyta realybė, leidžianti „pasimatuoti“ drabužius skaitmeninėje erdvėje. Dirbtinio intelekto įrankiai įgalina negalią turinčius asmenis per vaizdinio turinio įgarsinimą, balsu valdomus pokalbių robotus.

Apsiribojant vaizdinio ir tekstinio turinio kūrimu, svarbu suprasti kaip DI turinys yra kuriamas bei kokie veiksniai daro didžiausią įtaką tokio turinio kokybei. DI turinio kūrimo įrankių skaitmeninėje erdvėje itin gausu: „Perplexity“, „Gemini“, „DeepSeek“, „Runway“, „Canva“, „Pixlr Express“ ir „ChatGPT“. Pastarieji DI sprendimai vartotojams yra pasiekiami be jokių finansinių investicijų ir suteikia ribotą prieigą prie didžiosios dalies funkcijų.

Šiuolaikinio DI teksto generavimo pagrindą sudaro transformerių architektūros (angl. *transformer-based architectures*), pavyzdžiui tokie modeliai kaip GPT-4 (angl. *Generative Pre-trained Transformer 4*). Šie įrankiai gali kurti sklandų ir į kontekstą orientuotą tekstą (Brown et al., 2020). DI įrankiai „Taplio“ ir „OwlyWriter AI“ automatizuoja „LinkedIn“ įrašų kūrimą, perrašo jau egzistuojantį turinį ir net generuoja idėjas, kurios turi potencialą tapti populiariomis. DI modeliai yra apmokomi su didžiuliais tekstų kiekiais, todėl jie gali imituoti skirtingus rašymo stilius ir pritaikyti turinį konkrečioms platformoms, pavyzdžiui, „LinkedIn“ ar „Instagram“ (Brown et al., 2020). Siekiant sukurti profesionalų turinį DI pagalba, itin svarbi užklauskos formuluotė. Žemiau esančioje lentelėje (žr. 8 lentelę) pateikiamas užduoties pavyzdys bei išskiriamos struktūrinės jos dalys.

8 lentelė. Užklauskos generatyviniam dirbtiniam intelektui struktūra ir pavyzdys (sukurta autoriaus remiantis IBM, 2025b)

Struktūrinė dalis	Pavyzdys	Pavyzdys lietuvių kalba
-------------------	----------	-------------------------

Užduotis		Generate a LinkedIn post about quantum computing advancements in 2024	Sukurk įrašą „LinkedIn“ apie kvantinės kompiuterijos proveržius 2024-aisiais.
Gairės	Kalbėjimo tonas	Use professional yet accessible tone	Naudokite profesionalų, tačiau prieinamą toną
	Struktūra	Structure: 1. Hook with surprising statistic 2. Explain one breakthrough technology 3. Discuss industry implications 4. End with open question to drive engagement	Struktūra: 1. Pradėk nuo statistikos 2. Aprašyk vieną technologinį proveržį 3. Aptark pasekmes pramonei 4. Užbaik atviru, įsitraukimą skatinančiu klausimu
	Ilgis	180-220 words	180-220 žodžių
	Grotazymės	#QuantumComputing #TechInnovation	#KvantinėKompiuterija #TechnologėsInovacijos

Rekomenduotina užklausas formuluoti angliškai, nurodyti ne tik užduoties esmę, bet ir kalbėjimo toną, preliminarą struktūrą, ilgį, platformą, kurioje planuojama turinį skelbti (IBM, 2025b).

Kai kurie DI įrankiai, pavyzdžiui, „GoCharlie“ ar „FeedHive“, specializuojasi ilgo formato tekstų perdurbime ir pritaiko juos skirtingoms platformoms. Pavyzdžiui, jie gali iš tinklaraščio įrašo išskirti pagrindines mintis ir jas pritaikyti trumpam „Twitter“ įrašui ar „TikTok“ vaizdo įrašų antraštėms. Tam pasitelkiami pažangūs DI modeliai, tokie kaip BERT, kurie gali tinkamai suprasti ir interpretuoti teksto prasmę ir pritaikyti turinį taip, kad būtų išlaikyta pagrindinė mintis. Tokie dirbtinio intelekto įrankiai kaip „Replai.so“, leidžia generuoti natūralius ir į kontekstą orientuotus atsakymus į kitų socialinių tinklų vartotojų komentarus.

Vaizdinio turinio generavimas daugiausia remiasi generatyviais prieštariniais tinklais (angl. *Generative Adversarial Networks*, GANs) ir difuziniais modeliais. GANs, kuriuos pristatė Goodfellow‘as et al. (2014), veikimas apibūdinamas kaip: du tarpusavyje konkuruojantys neuroniniai tinklai – generatorius (angl. *generator*) ir diskriminatorius (angl. *discriminator*) – kurių tarpusavio sąveika sukuria realistiškus vaizdus. Generatorius kuria netikrus vaizdus ir siunčia juos diskriminatoriui. Diskriminatorius bando atpažinti, ar gautas pavyzdys yra tikras, ar sukurtas DI. Jei diskriminatorius teisingai atpažįsta sintetinius duomenis, generatorius gauna neigiamą grįžtamąjį ryšį ir mokosi kurti tikroviškesnius vaizdus. Jei diskriminatorius suklysta ir palaiko sugeneruotą duomenį tikru, generatorius gauna teigiamą grįžtamąjį ryšį ir tęsia tokį mokymąsi. Laikui bėgant abu tinklai tampa vis stipresni – generatorius išmoka kurti vis realistiškesnį turinį, o diskriminatorius geba geriau jį atpažinti. Tokios priemonės kaip „Publer“ ir „BannerBear“ naudoja GANs variantus nuotraukoms, logotipams ir reklaminėms vizualizacijoms kurti. Tuo tarpu difuziniai modeliai, tokie kaip „Stable Diffusion“, leidžia vartotojams sukurti aukštos kokybės vaizdus pagal teksto užklausas. Visgi, Rombach‘as et al. (2021) pabrėžia, kad difuziniai vaizdų generavimo modeliai reikalauja daug kompiuterinių išteklių ir yra sunkiai pritaikomi, jei sugeneruoto turinio reikia čia ir dabar. Užduoties formuluotė vaizdai generuoti šiek tiek skiriasi nuo teksto. Šįkart itin svarbu apibrėžti tiek norimą vaizdą, tiek technines vaizdo specifikacijas, pavyzdžiui, „*a futuristic cityscape at twilight, 8k resolution, cinematic lighting, neo-cyberpunk architecture with glowing neon accents, hyper-detailed textures, unreal engine 5 rendering, depth of field*“ (IBM, 2025b).

Modernūs socialinių tinklų valdymo įrankiai, tokie kaip „FeedHive“ ar „Ocoya“, sujungia tekstų ir vaizdų generavimą, kad palengvintų turinio kūrimą. Tam naudojami daugiamodaliniai DI modeliai, tokie kaip CLIP (angl. *Contrastive Language–Image Pretraining*), kurie gali suderinti tekstinį ir vaizdinį turinį (Radford et al., 2021).

Egzistuojančių GPT modelių modifikavimas pagal konkrečius poreikius, gali padidinti jų efektyvumą ir aktualumą. Adaptuoti DI teksto ir vaizdo generavimo modeliai į generuojamus duomenis įtraukia konkrečios srities žinias, terminologiją ir kontekstą, todėl pateikiami rezultatai tampa tikslesni ir naudingesni. Pavyzdžiui, „AcademicGPT“ – modelis, pritaikytas akademiniams tekstams, buvo sukurtas siekiant padėti tyrėjams atsakyti į akademinius klausimus, apibendrinti mokslinius straipsnius ir generuoti tyrimų idėjas. Ši specializacija leidžia tyrėjams efektyviau analizuoti ir apdoroti didžiulį mokslinės informacijos kiekį (Wei et al., 2023).

DI generuojamas turinys keičia socialinius tinklus – jis leidžia kurti efektyvesnį ir greitesnį turinį, tačiau kartu kelia ir naujų iššūkių. Svarbu rasti pusiausvyrą tarp automatizacijos ir žmogaus įsikišimo, kad sukurtas turinys būtų etiškas, patikimas ir naudingas vartotojams.

2.1.3. Žmogaus kuriamo turinio socialiniuose tinkluose charakteristikos

Socialiniai tinklai tapo reikšminga platforma, kurioje atsiskleidžia unikalūs žmogaus kūrybinis indėlis. Ši išskirtinumo samprata remiasi ne tik techniniais ar vizualiniais kūrybos aspektais, bet ir autentiškais žmogaus išgyvenimais, emocijomis bei vertybėmis, kurios perteikiamos per kuriamą turinį (Kemp et al., 2021). Skirtingai nei dirbtinio intelekto generuojamas turinys, kuris dažnai pasižymi standartizuota ir numanoma logika (Berber Sardinha, 2024), žmogaus kūryba pasižymi spontaniškumu, emociniu gilumu ir gebėjimu perteikti subtilius kultūrinius niuansus (Kemp et al., 2021).

Žmogaus kuriamam turiniui būdingos kelios esminės charakteristikos, kurios pabrėžia jo unikalumą ir reikšmę, siekiant kurti efektyvią ir įtraukiančią komunikaciją. Toliau pateikiamoje lentelėje (žr. 8 lentelę) apibendrintos pagrindinės šios savybės, jų reikšmė (Belanche et al., 2024; Berber Sardinha, 2024; Guan et al., 2022; Hofacker et al., 2020; Ji et al., 2023; Kemp et al., 2021; Krishen et al., 2021). Siekiant iliustruoti charakteristikas pateikiami pavyzdžiai iš populiarių prekių ženklų praktikos .

9 lentelė. Žmogaus kuriamo turinio charakteristikos

Charakteristika		Reikšmė	Pavyzdys
Personalizacija ir empatija	Kontekstinis supratimas	Turinys pritaikomas prie kultūrinio, socialinio ir situacinio konteksto.	„McDonald's“: skirtingose šalyse kuria lokalizuotus meniu ir reklamines kampanijas, pavyzdžiui, „McAloo Tikki“ Indijoje.
	Empatija	Gebėjimas atpažinti auditorijos emocinius poreikius ir kurti ryšį.	„Dove“: „Real Beauty“ kampanija, akcentuojanti tikrų moterų grožį ir stiprinanti savivertę.
	Individualizacija	Turinys pritaikomas konkrečioms vartotojams pagal jų pomėgius ir poreikius.	„Spotify“: „Wrapped“ metų apžvalga, kuri personalizuoja vartotojų muzikos klausymo statistiką ir dalijimosi galimybę.

Charakteristika		Reikšmė	Pavyzdys
Kūrybiškumas ir emocinis poveikis	Originalumas	Turinys kyla iš asmeninės patirties, emocinio supratimo ir tikro požiūrio.	Realios klientų istorijos, pavyzdžiui, „Nike“ sportininkų pasakojimai apie sėkmę.
	Istorijų pasakojimas	Įkvepiančios ir emocionalių įtraukiančios istorijos.	„Coca-Cola“ kampanijos apie draugystę ar šeimą.
	Emocinis poveikis	Turinys, kuris sukelia stiprias emocijas ir skatina veiksmą.	„Always“: „Like a Girl“ kampanija, pabrėžianti moterų įgalinimą ir stereotipų laužymą.
Lankstumas ir adaptacija	Kūrybinis lankstumas	Gebėjimas prisitaikyti prie naujų tendencijų ir rinkos pokyčių.	„TikTok“: kampanijos, pasitelkiančios virusinius iššūkius, pavyzdžiui, „Guess“ #InMyDenim iššūkis.
	Socialinis poveikis	Turinys, kuris skatina dialogą ir bendruomenės jausmą tarp vartotojų.	„LEGO“: „Rebuild the World“ kampanija, skatinanti vartotojus dalytis kūrybinėmis idėjomis ir kurti bendruomenę.

Vienas svarbiausių žmogaus kuriamo turinio bruožų yra jo gebėjimas personalizuoti žinutę, pritaikant ją konkrečiai auditorijai ir atsižvelgiant į kontekstą. Guan et al. (2022) tyrime atskleidžiama, kad žmogaus sukurtas turinys dažnai atspindi kultūrinius ir socialinius skirtumus, kurie stiprina turinio aktualumą ir poveikį specifinėms auditorijoms. Personalizacija žmogaus kūryboje kyla iš empatijos, intuicijos ir gebėjimo atpažinti subtilius auditorijos poreikius bei nuotaikas. Tai gali būti siejama su žmogaus gebėjimu reaguoti į specifinius įvykius ar situacijas – kultūrinį šventinį kontekstą ar visuomenės nuotaikų pokyčius. Pavyzdžiui, „McDonald's“ savo reklamas ir meniu pritaiko skirtingose šalyse, įtraukdami vietinius skonius ir tradicijas, kaip Indijos „McAloo Tikki“, kuris atitinka vietines vegetarizmo tradicijas. Skirtingai nei dirbtinio intelekto pagalba grindžiama personalizacija, paremta dideliu duomenų kiekiu, algoritmais ir iš anksto nustatytais modeliais prognozėms atlikti, žmogiškoji personalizacija kur kas gilesnė ir labiau intuityvi. DI analizuoja vartotojų elgseną, remdamasis jų buvusią veiklą – paspaudimais, pirkiniais ar naršymo istorija, – ir, algoritmų pagalba, apdoroja šią informaciją, siekiant pasiūlyti suasmenintus sprendimus (Berber Sardinha, 2023). DI gali pritaikyti turinį pagal vartotojo anksčiau peržiūrėtas prekes ar aplankytus internetinius puslapius, tačiau jam trūksta žmogiškojo gebėjimo įvertinti platesnį kontekstą, tokį kaip vartotojo emocinė būsena ar subtilūs kultūriniai niuansai. Personalizacija, kurią atlieka žmogus, priešingai, remiasi ne tik tuo, kas buvo padaryta praeityje, bet ir gebėjimu įvertinti dabartinę situaciją bei prognozuoti auditorijos emocinę reakciją į tam tikrą turinį. Žmogaus kūrybinis indėlis kuria ne tik emocinį ryšį, bet ir ilgalaikį pasitikėjimą. Kaip pabrėžia Berber'is Sardinha's (2024), personalizuota komunikacija, pasižyminti nuoširdumu ir dėmesiu auditorijos poreikiams, yra itin vertinama ir padeda formuoti lojalią bendruomenę, o vartotojai jaučiasi išgirsti ir suprasti.

Žmogaus kuriamame turinyje išryškėja empatijos ir kultūrinio konteksto svarba. Guan et al. (2022) atskleidė, kad skirtingas kultūrinis kontekstas ne tik formuoja auditorijos lūkesčius ir požiūrį, bet ir daro įtaką tam, kokio pobūdžio turinys yra laikomas vertingu. Šis suvokimas leidžia turinio kūrėjams geriau prisitaikyti prie specifinių auditorijų poreikių, kuriant turinį, kuris yra tiek emocingas, tiek informatyvus. Kolektyvistinėse kultūrose, pavyzdžiui, Rytų Azijos šalyse, didelis dėmesys skiriamas tarpusavio santykiams ir bendruomeniškumui, empatija ir gilūs emociniai aspektai yra pagrindinės vertybės, kurios skatina vartotojų įsitraukimą. Šių regionų auditorijos vertina turinį, kuris atspindi

harmoniją, bendruomenės vertybes ir emocinį gilumą, nes tai atliepia jų kasdienius išgyvenimus ir pasaulėžiūrą. Tuo tarpu individualistinėse kultūrose, tokiose kaip Jungtinės Amerikos Valstijos, turinys, kuris pabrėžia asmeninius pasiekimus, saviraišką ir unikalumą, yra suvokiamas kaip itin patrauklus. Empatija šiuo atveju išreiškiama per individualius pasakojimus ir autentišką kūrėjų balsą, kuris leidžia vartotojams jaustis asmeniškai įtrauktiems į istoriją. Pavyzdžiui, „Dove“ kampanija „Real Beauty“ pabrėžia autentišką moterų grožį ir vertybes, stiprindama emocinį ryšį su auditorija, o tai padeda kurti ilgalaikį lojalumą, nes vartotojai jaučiasi svarbia bendruomenės dalimi (Ji et al., 2023). Kultūrinis kontekstas gali paveikti net ir turinio formatą bei stilių: kai kuriose kultūrose vaizdo įrašai ar vizualiniai pasakojimai labiau atitinka auditorijos lūkesčius, nes jie greičiau perteikia emocinį turinį. Kitur, kaip, pavyzdžiui, Europoje, detalūs tekstiniai pasakojimai ar apmąstymai yra laikomi vertingesniais, nes jie sukuria gilesnį intelektualinį ir emocinį ryšį. Empatijos svarba ypač atsiskleidžia komunikuojant jautriomis temomis, pavyzdžiui, krizių metu. Šiuo atveju auditorijos iš kultūrų, turinčių stiprų kolektyvinį identitetą, tokiais atvejais tikisi, kad turinys perteiks ne tik informaciją, bet ir palaikymo žinutę, kuri stiprina solidarumo jausmą. Tai reiškia, kad turinio kūrėjai turi ne tik informuoti, bet ir tapti emocinio palaikymo šaltiniu, kuris rezonuoja su specifiniais auditorijos poreikiais (Guan et al., 2022; Krishen et al., 2021).

Belanche et al. (2024) pabrėžia, kad žmogaus sukurtas turinys išsiskiria gebėjimu užmegzti artimą ryšį su auditorija per identifikaciją ir empatiją, kurių dirbtiniam intelektui dažnai trūksta. Realūs nuomonės formuotojai, dalindamiesi asmenine patirtimi, emocijomis ir gyvenimo akimirkomis, sukuria tvirtą socialinį ryšį su sekėjais. Pastarasis ypač svarbus reklamuojant hedoniškus produktus, tokius kaip prabangūs viešbučiai ar mados prekės. Ši artima sąveika leidžia vartotojams geriau susitapatinti su nuomonės formuotoju ir patirti emocinį pasitenkinimą, kuris padidina norą sekti pateikiamomis rekomendacijomis. Tuo tarpu dirbtinio intelekto pagrindu kuriami virtualūs nuomonės formuotojai yra labiau pritaikyti utilitarinių produktų, tokių kaip technologiniai sprendimai ar buitiniai prietaisai, reklamai. Virtualūs nuomonės formuotojai pasižymi objektyvumu ir analitiniu požiūriu, kuris suteikia jų komunikacijai didesnę naudingumo pojūtį. Ši savybė yra itin vertinama vartotojų, ieškančių logiškai pagrįstų ir informatyvių rekomendacijų.

Žmogaus kuriamas turinys įtraukia auditoriją pasiekdamas stiprų emocinį poveikį per gebėjimą pasakoti įtraukiančias istorijas. Kemp et al. (2021) pabrėžia istorijų pasakojimo (angl. *storytelling*) svarbą, kuri ne tik leidžia kurti ryšius tarp prekių ženklų ir vartotojų, bet ir padeda įtraukti auditoriją į turinio kūrimo procesą. Istorijų pasakojimas yra galinga priemonė, leidžianti vartotojams tapti pasakojimo dalimi ir asmeniškai susitapatinti su prekės ženklu. Tai sukuria emocinį ryšį, kuris ne tik didina auditorijos įsitraukimą, bet ir ilgalaikį lojalumą. Vienas iš svarbiausių istorijų pasakojimo elementų yra emocinis transportavimas – procesas, kurio metu auditorija pasineria į pasakojimą, identifikuojasi su personažais ar situacijomis ir patiria stiprias emocijas (Kemp et al., 2021). Toks asmeninio ryšio su auditorija sukūrimas leidžia skatinti vartotojus aktyviai dalyvauti turinio sklaidoje. Hofacker'is et al. (2020) pažymi, kad vartotojų dalyvavimas kuriant ir dalinantis turiniu, įkvėptu istorijų, stiprina prekės ženklo patikimumą ir kuria bendruomenės jausmą. Svarbu tai, kad istorijų pasakojimas leidžia ne tik perduoti informaciją, bet ir daryt įtaką auditorijos elgsenai. Kemp et al. (2021) išskiria, kad gerai struktūruotos istorijos, kuriose pabrėžiami herojiški veikėjų poelgiai ar jų sėkmės istorijos, ne tik įtraukia auditoriją, bet ir skatina ją veikti. Pavyzdžiui, prekių ženklai, tokie kaip „Nike“ ar „Coca-Cola“ dažnai naudoja realių klientų sėkmės istorijas, kurios tampa įkvėpimu kitai auditorijos daliai. Tokios istorijos sustiprina vartotojų pasitikėjimą prekės ženklu, bet ir padeda formuoti jų savivoką bei norą būti dalimi šios istorijos. Tuo tarpu dirbtiniu intelektu grįstiems

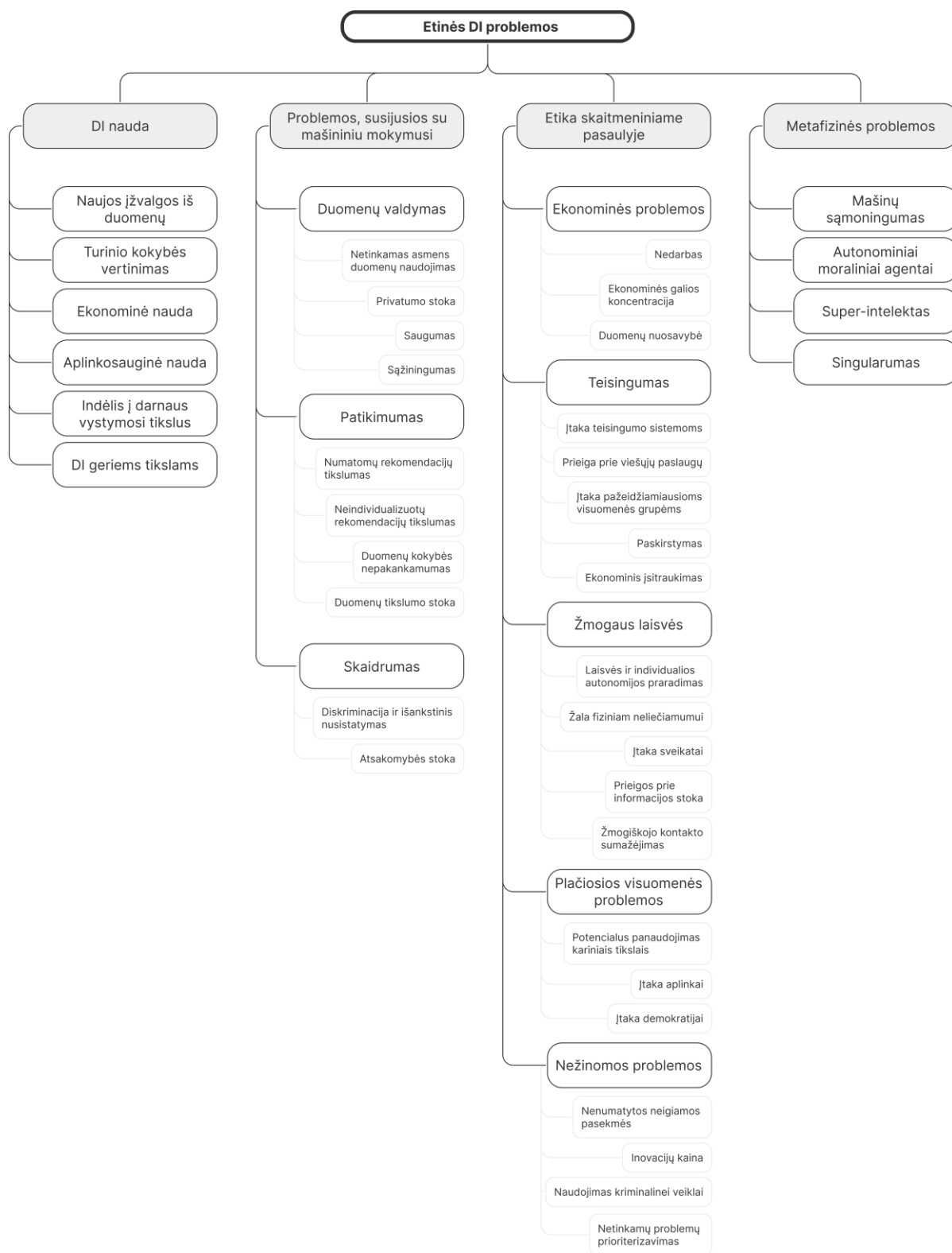
sprendimams neretai sunku kurti įtikinamas istorijas, nes jis neturi asmeninės patirties, emocinio supratimo ir gebėjimo įsijausti, o tai būtina autentiško ryšio su auditorija kūrimui. Be to, DI dažnai stokoja subtilaus kultūrinio ir socialinio konteksto suvokimo, kuris yra esminis efektyvioms istorijoms pasakoti. Guan et al. (2022) atkreipia dėmesį, kad vartotojų skatinimas kurti savo istorijas ar dalintis išpūdžiais apie produktus socialiniuose tinkluose tampa svarbia integruotos rinkodaros dalimi. Tokiu būdu prekės ženklas ne tik įtraukia vartotojus į kūrybinį procesą, bet ir stiprina jų asmeninį ryšį su turiniu bei skatina natūralią turinio sklaidą.

Apibendrinant, žmogaus kuriamas turinys socialiniuose tinkluose pasižymi autentiškumu, emociniu gilumu ir gebėjimu prisitaikyti prie įvairių auditorijų poreikių. Pagrindinės šio turinio savybės, tokios kaip personalizacija, empatija, autentiškumas, istorijų pasakojimas ir kūrybinis lankstumas, leidžia užmegzti stiprų emocinį ryšį su vartotojais bei formuoti jų lojalumą. Skirtingai nuo dirbtinio intelekto generuojamo turinio, grindžiamo standartizuotomis prognozėmis ir pasižyminčio emocinio jautrumo stygiumi, žmogaus kuriamas turinys išlaiko unikalumą ir intuityvumą. Tokiu būdu žmogaus kūryba tampa pagrindine priemone ne tik pasiekti auditoriją, bet ir ją įtraukti į turinio kūrimo bei sklaidos procesą, kurti vertę, kuri ilgainiui stiprina ryšį tarp vartotojų ir prekių ženklų.

2.2. Vartotojų pasitikėjimas dirbtiniu intelektu ir jo etiniai aspektai

Dirbtinis intelektas tampa vis svarbesniu įrankiu įvairiose srityse, nuo sveikatos priežiūros ir finansinių technologijų iki švietimo ir socialinių tinklų, tačiau jo priėmimas ir naudojimas didžiąja dalimi priklauso nuo vartotojų pasitikėjimo. Afroogh'as et al. (2024) pabrėžia, kad pasitikėjimas DI nėra vien tik techninis aspektas, susijęs su sistemos tikslumu ar patikimumu. Tai taip pat apima ir aksiologinę dimensiją, kurioje pagrindinis dėmesys skiriamas etinėms ir socialinėms vertybėms, tokioms kaip atsakingas naudojimas, sprendimų paaiškinamumas ir galimos žalos mažinimas. Pavyzdžiui, DI technologijos gali kelti grėsmių vartotojų autonomijai, jų privatumui ar net duomenų saugumui, jei šie aspektai nėra tinkamai valdomi. Skaidrumo trūkumas yra viena didžiausių kliūčių – vartotojai dažnai neturi aiškaus supratimo, kaip DI priima sprendimus, o tai gali kelti nepasitikėjimą technologija. Felzmann et al. (2020) pažymi, kad skaidrumo didinimas yra esminė sąlyga pasitikėjimui kurti, tačiau jis turi būti įgyvendintas taip, kad vartotojai galėtų lengvai suprasti DI veikimą, net jei jie neturi techninių žinių. Dar viena problema yra dezinformacijos plitimas per DI generuotą turinį. Nors DI sprendimai itin reikšmingai pasitarnauja siekiant mažinti dezinformacijos sklaidą, jis taip pat gali prisidėti prie dezinformacijos sklaidos, generuodamas klaidingus tekstus ar vaizdus (Lu et al., 2022). Svarbu pabrėžti, jog pasitikėjimo ir rizikos balansas yra esminis veiksnys, ypač jautrioje srityse, tokiose kaip sveikatos apsauga ar finansai. Per didelį pasitikėjimą DI gali lemti vartotojų negebėjimą atpažinti sistemos klaidas, o vartotojų nepasitikėjimas technologija gali mažinti sprendimų naudingumą ir priimtinumą.

Stahl'as et al. (2021) išskyrė dirbtinio intelekto naudas bei tris kategorijas, apimančias pagrindines etines problemas DI kontekste. Dauguma iš išskirtų temų yra detalizuojamos į smulkesnes subkategorijas (žr. 1 pav.).



1 pav. Etinės DI problemos (pagal Stahl et al., 2021)

Autoriai išskiria dirbtinio intelekto privalumus, tokius kaip efektyvumo didinimas, ekonominė nauda ir indėlis į darnaus vystymosi tikslus, tačiau taip pat pažymi ir rimtus etinius iššūkius. Paveiksle išskirtos trys pagrindinės etinių problemų kategorijos: problemas, kylančias iš mašininio mokymosi (pavyzdžiui, duomenų privatumo pažeidimai, šališkumas ir skaidrumo trūkumas), DI poveikis visuomenei (pavyzdžiui, darbo rinkos pokyčiai, ekonominės galios koncentracija, žmogaus autonomijos mažėjimas) ir metafizines problemas (pavyzdžiui, mašinų sąmoningumas ar

superintelektas). Straipsnio autoriai pabrėžia, kad nors organizacijos yra sąmoningos ir atkreipia dėmesį į dalį iš šių iššūkių, jų taikomos priemonės, tokios kaip duomenų anonimizavimas ir DI sprendimų priežiūra, neapima visų iškylančių problemų. DI etikos užtikrinimas reikalauja ne tik iniciatyvų, kylančių organizacijų viduje, bet ir plataus masto reguliavimo bei politikų pokyčių.

Paaškinamumas, arba gebėjimas suprasti kaip DI sistemos priima sprendimus, yra kertinė dedamoji, siekiant ugdyti pasitikėjimą ir užtikrinti atskaitomybę (Singhal et al., 2024a). DI algoritmai socialinių tinklų platformose dažnai veikia kaip „juodosios dėžės“, todėl sunku suprasti, kodėl tam tikras turinys yra akcentuojamas ar iškeliamas (Joy et al., 2024). Tokios technikos kaip LIME (angl. *Local Interpretable Model-agnostic Explanations*) ir SHAP (angl. *SHapley Additive exPlanations*) gali padėti geriau suprasti šiuos procesus, tačiau jos nepajėgia visiškai išspręsti etinio skaidrumo poreikio (Singhal et al., 2024a). Dabartiniai paaškinamumo metodai dažniausiai teikia tik aprašomąsias įžvalgas, kurios apibūdina kaip buvo priimtas sprendimas, o ne normatyvius sprendimų pagrindimus, atsakančius į klausimą, ar šis sprendimas yra teisingas ir pagrįstas pagal moralinius, teisinius ar socialinius standartus. Tai pabrėžia žmogaus atsakomybės svarbą kuriant ir diegiant dirbtinį intelektą. Joy et al., (2024) tikina, kad socialinių tinklų platformos privalo teikti pirmenybę skaidrumui, viešai atskleidžiamos, kaip veikia jų algoritmai ir kokiais kriterijais remiamasi moderuojant turinį.

Atsakomybės klausimas DI kontekste yra sudėtingas reiškinys, kadangi jis įtraukia keletą skirtingų šalių: įrankio kūrėjus (angl. *developers*), diegėjus (angl. *deployers*) ir naudotojus (angl. *users*). Nereguliuojamos DI sistemos socialiniuose tinkluose gali būti kaip terpė dezinformacijos sklaidai, neigiamo nusistatymo stiprinimui ar asmens duomenų viešinimui, o tai, savo ruožtu, didina poreikį nustatyti, kas galėtų būti atsakingas už šiuos nusižengimus. Loconsole ir Panarari's (2025) tikina, kad atsakomybės klausimas DI kontekste turi būti įvairialypis, sukoncentruotas tiek į techninius, tiek į organizacinius aspektus. Techniniu lygmeniu, paaškinamumas padeda diagnozuoti ir įvertinti trikdžius DI sprendimų priėmimo procese. Tuo tarpu organizaciniame lygmenyje gali būti taikomos teisinės ir reguliacinės priemonės, pavyzdžiui, Europos Sąjungos dirbtinio intelekto akte (angl. *ES AI Act*) (European Union, 2024) pateikiama rizikos pagrindu grindžiama DI sistemų klasifikacija, pagal kurią atsakomybė už vienas ar kitas dirbtinio intelekto klaidas yra priskiriama ir žmogui. Pagal šį aktą, atsakomybę už DI sistemų klaidas tenka skirtingiems subjektams (žr. 10 lentelę), priklausomai nuo jų vaidmens kuriant, diegiant ar naudojant DI sistemą.

10 lentelė. Atsakomybė už DI klaidas pagal subjektą (sudaryta remiantis ES dirbtinio intelekto aktu)

Subjektas	Pavyzdys
Kūrėjai (angl. <i>developers</i>)	Jei DI sistema, pavyzdžiui, autonominis automobilis, sukelia avariją dėl programinės įrangos klaidos, atsakomybė gali tekti sistemos kūrėjui. Tokiu atveju nukentėję asmenys turi teisę reikalauti kompensacijos iš tiekėjo už patirtą žalą.
Diegėjai (angl. <i>deployers</i>)	Jeigu bankas naudoja DI pagrįstą kredito vertinimo sistemą, kuri dėl netinkamo diegimo ar priežiūros klaidingai atmeta klientų paraiškas, atsakomybė gali tekti bankui kaip sistemos diegėjui. Tokiu atveju klientai gali reikalauti kompensacijos už neteisingą sprendimą.
Naudotojai (angl. <i>users</i>)	Jei asmuo naudoja DI sistemą neteisėtais tikslais arba pažeisdamas naudojimo sąlygas, atsakomybė už padarytą žalą gali tekti pačiam vartotojui. Pavyzdžiui, jei asmuo naudoja DI įrankį dezinformacijai skleisti, jis gali būti laikomas atsakingu už sukeltą žalą.

Europos Sąjungos DI akte numatytos griežtos baudos už dirbtinio intelekto akto pažeidimus. Pavyzdžiui, didelėms įmonėms už draudžiamų DI praktikų naudojimą gali būti skiriamos baudos iki 35 mln. eurų arba iki 7 proc. metinės pasaulinės apyvartos, priklausomai nuo to, kuri suma yra didesnė. Tačiau mažoms, vidutinio dydžio įmonėms bei startuoliams numatytos pagal pajamas

proporcingai mažesnės administracinių baudų ribos. Socialinių medijų platformų valdytojai privalo neapsiriboti vien tik reguliavimo atitiktimi – jos turėtų aktyviai kurti aiškas etikos gaires, vykdyti nepriklausomus auditus ir steigti etikos komitetus, siekdamos užtikrinti, kad DI vystymas ir taikymas būtų suderintas su visuomenės vertybėmis (Loconsole ir Panarari, 2025).

Dirbtinis intelektas atlieka prieštarinę vaidmenį dezinformacijos ekosistemoje – jis vienu metu prisideda prie klaidingos informacijos sklaidos ir padeda ją suvaldyti (Monteith et al., 2024). Dezinformacijos stiprinimas DI pagrindu veikiančiose sistemose kyla iš su įsitraukimu susijusios optimizacijos, kur rekomendavimo algoritmai dažniausiai prioritetą teikia turiniui, kuris skatina didžiausią vartotojų įsitraukimą, o ne tam, kuris yra tiksliausias. Dėl šios priežasties emociniu požiūriu stiprus ir sensacingas turinys plinta greičiau nei patikima informacija (Yildiz Durak et al., 2025). Be to, DI modeliai, tokie kaip GPT-4, gali itin greitai generuoti įtikinamą, nuoseklią ir į kontekstą orientuotą dezinformaciją, o tai, savo ruožtu, leidžia piktavaliams kurti netikras naujienas, manipuliuoti viešąja nuomone ir skleisti melagingas žinutes socialinių medijų platformose. Tuo pat metu DI sistemos naudojamos, siekiant optimizuoti savo dezinformacijos sklaidą, adaptuojant klaidinančias žinutes pagal vartotojų poreikius. Tai sukuria sudėtingą etinį iššūkį – tos pačios technologijos, kurios leidžia efektyviai kurti ir platinti turinį, taip pat tampa priemone klaidinančiai informacijai skleisti (Monteith et al., 2024). Kita vertus, DI sprendimai taip pat pasitelkiami dezinformacijos aptikimui ir neutralizavimui. Automatizuotos faktų tikrinimo sistemos pasitelkia natūralios kalbos apdorojimo (angl. *Natural language processing, NLP*) metodus, kad palygintų teiginius su patikimais šaltiniais, o kompiuterinės regos technologijos (angl. *Computer vision technologies*) gali nustatyti vaizdo klastotes (angl. *deepfakes*). Tuo tarpu dirbtiniu intelektu paremta tinklų analizė padeda identifikuoti koordinuotus dezinformacijos platinimo veiksmus (X. Chen et al., 2022). Tačiau skaidrumas, kaip viena iš dezinformacijos kontrolės priemonių, kelia papildomų iššūkių. Paaiškinamas DI (angl. *Explainable AI, XAI*) gali informuoti kenkėjus, kodėl tam tikras turinys buvo pažymėtas (angl. *flagged*) kaip klaidinantis ar netinkamas. Tuo pasinaudojus, įmanoma sukurti sistemas, pajėgias apeiti automatizuotas melagingo ar netinkamo turinio aptikimo sistemas (Monteith et al., 2024). Siekiant efektyviai kontroliuoti dezinformacijos plitimą, būtinas įvairiapusis požiūris: turinys turėtų būti reitinguojamas ne tik pagal populiarumą, bet ir pagal tikslumą; svarbu naudoti sistemas, kurios gali realiu laiku tikrinti faktus ir nustatyti, ar informacija yra tiksli; galiausiai, turinio moderavimą turėtų atlikti ne tik DI algoritmai, bet ir žmonės. Ne ką mažiau svarbus ir pačių vartotojų socialinės žiniasklaidos raštingumas bei kritinis mąstymas – pagrindiniai įrankiai, padėsiantys atpažinti melagingą informaciją (McDermid et al., 2021).

Vartotojų švietimas apie DI vaidmenį socialinėje medijoje taip pat yra labai svarbus siekiant sukurti informuotą ir atsparią skaitmeninę visuomenę. Daugelis žmonių nesuvokia, koku mastu DI sistemos formuoja jų internetinę patirtį – nuo naujienų srautų sudarymo iki komentarų filtravimo ar dezinformacijos atpažinimo (Joy et al., 2024). Neturėdami šių žinių, vartotojai gali pasyviai priimti algoritmų sprendimus, nekvestionuodami jų tikslumo, teisingumo ar galimų šališkumų. Supratimas apie DI įtaką padeda žmonėms kritiškai vertinti matomą turinį, sumažina manipuliacijos riziką ir didina sąmoningumą apie informacijos patikimumą (Gilbert et al., 2020). Be to, žinios apie DI ribotumus – pavyzdžiui, tai, kad jis ne visada geba atskirti satyrą nuo tikrų naujienų – gali padėti vartotojams nepasikliauti DI sprendimais akiai (Niu Yongbo, 2025). Singhal'as et al. (2024) pažymi, kad svarbu investuoti į žiniasklaidos raštingumo programas, suteikiančias vartotojams prieinamus išteklius, paaiškinančius, kaip DI kuria turinį ir kaip galima koreguoti nustatymus, kad patirtis būtų kuo labiau individualizuota. Tokios funkcijos kaip pritaikomos rekomendacijų sistemos, skaidrumo

informaciniai skydeliai ir aiškūs turinio moderavimo paaiškinimai gali suteikti vartotojams daugiau galimybių kontroliuoti savo internetinę patirtį. Žmonės, turėdami šias priemones, yra skatinami tapti savarankiškais skaitmeninėje erdvėje, taip ilgainiui sukuriant patikimesnę ir atsakingesnę DI ekosistemą. (Joy et al., 2024). Kadangi DI algoritmai gali sustiprinti esamus šališkumus ir sukelti nepageidaujamų padarinių, būtina užtikrinti, kad vartotojai turėtų priemones atpažinti šias problemas ir aktyviai dalyvauti diskusijose apie DI valdymą (Gilbert et al., 2020).

Apibendrinant, galima teigti, kad šiuolaikinėje skaitmeninėje erdvėje, įskaitant ir socialinius tinklus, dirbtinio intelekto etinių aspektų valdymas reikalauja visapusiško ir bendradarbiavimu grindžiamo požiūrio, įtraukiančio vyriausybes, pramonės lyderius, mokslininkus ir pačią visuomenę. Norint užtikrinti vartotojų pasitikėjimą DI sistemomis, būtina užtikrinti jų skaidrumą, sąžiningumą, privatumą ir naudotojų autonomiją. Itin svarbus yra ir pačio vartotojo – žmogaus – vaidmuo – tiek technologijų kūrėjai, tiek naudotojai turi užtikrinti etišką DI taikymą savo veikloje.

2.3. Turinio socialiniuose tinkluose autentiškumo suvokimas

Pastaraisiais metais auganti dirbtinio intelekto technologijų įtaka formuoja naujas turinio kūrimo ir vertinimo normas, kuriose vartotojų gebėjimas nustatyti turinio autentiškumą tampa kritiškai svarbus. Maares et al. (2021) pabrėžia, kad socialiniuose tinkluose autentiškumas apima ne tik tikrumą ar tiesą, bet ir gebėjimą perteikti vertybes bei emocinį ryšį, kuris atitiktų vartotojų lūkesčius. Autentiškumas socialinėje žiniasklaidoje dažnai siejamas su sąžiningumu, natūralumu ir tikrumu. Jis gali būti vertinamas pagal tai, kaip vartotojai suvokia kūrėjų ar prekių ženklų veiklą. Lee (2020) nurodo, kad svarbiausi autentiškumo kriterijai yra skaidrumas, nuoseklumas ir tikėjimas pateikiama žinute. Tyrimai rodo, kad vartotojai vis labiau domisi ne tik turinio estetika, bet ir jo sąsaja su tikrosiomis vertybėmis. Jenkins et al. (2020) pastebi, kad net maži neatitikimai tarp žinutės ir jos perteikimo gali būti suvokiami kaip manipuliacija arba neautentiškumas, o tai lemia pasitikėjimo praradimą. Autentiškumas yra neatsiejamas nuo technologinių platformų galimybių ir iššūkių. Šiuolaikiniai algoritmai gali manipuluoti vartotojų suvokimu, skatindami tam tikrą turinį, tačiau tai dažnai kelia klausimų dėl tokio turinio tikrumo. J. Yang et al. (2021) atkreipia dėmesį, kad skirtingose socialinių tinklų platformose vartotojų lūkesčiai turinio autentiškumo atžvilgiu gali skirtis priklausomai nuo kultūrinių ir socialinių normų, tačiau nuoseklumas ir sąžiningumas išlieka pagrindiniais vertinimo kriterijais.

Autentiškumas socialiniuose tinkluose apibrėžiamas kaip turinio tikrumas ir atitikimas autoriaus tapatybei ar vertybėms. Kitaip tariant, turinys laikomas autentišku, kai matoma, jog kūrėjas „yra ištikimas sau“, – jo skelbiamos mintys ir elgesys atspindi tikrąją tapatybę (Jenkins et al., 2020b). Autorystės atskleidimas (pavyzdžiui, paviešinant tikrą vardą, profilį) suteikia kontekstą šiam autentiškumui įvertinti. Tyrimai rodo, jog žmonės linkę vertinti asmenį kaip autentišką, kai jo veiksmai suderinami su jo deklaruojama tapatybe. Pavyzdžiui, įžmybės, kurios laikomos autentiškomis, daro didesnę įtaką sekėjams tiek internete, tiek už jo ribų. Todėl skaidrus autorystės atskleidimas gali padėti vartotojams pajusti, kad turinys „tikras“ ir nesumeluotas (Kim et al., 2015).

Vis dėlto autentiškumo suvokimas labai priklauso nuo to, ar vartotojai žino, kas yra turinio autorius. Jei socialiniame tinkle pranešimas paskelbtas aiškiai įvardyto asmens ar organizacijos vardu, auditorija turi rėmus (pavyzdžiui, kiti organizacijos ar asmens įrašai, atsiliepimai apie juos forumuose ar reitingavimo platformose), per kuriuos interpretuoja žinutę. Priešingai, anoniminis turinys ar žinutės paskelbtos pseudonimu gali kelti klausimų dėl patikimumo – neaišku, kas slypi už žinutės, ir

kokie jos autoriaus motyvai. Socialiniuose tinkluose neretai taikomi heuristiniai vertinimai: vartotojai mažai įsigilinę į turinį vertina jį pagal lengvai pastebimas užuominas (angl. *peripheral cues*) (Jenkins et al., 2020b). Autorystė yra viena iš tokių užuominų – žinodami autorių, žmonės gali lengviau numanyti turinio patikimumą, vietoj detalios ir kruopščios pačios žinutės turinio analizės. Taigi, teoriniame lygmenyje autorystės atskleidimas sukuria sąlygas autentiškumui pasireikšti ir palengvina auditorijos sprendimą, siekiant nustatyti ar turinys vertas pasitikėjimo (Jenkins et al., 2020b; Kim et al., 2015).

Aiškiai nurodyta turinio autorystė stipriai veikia vartotojų pasitikėjimą turiniu. Pagal šaltinio patikimumo modelį (angl. *Source Credibility Model*), auditorijos pasitikėjimas turiniu labai priklauso nuo to, kaip vertinamas pats informacijos šaltinis pagal kompetenciją ir patikimumą (Seiler & Kucza, 2017). Jeigu vartotojai mano, kad žinutės autorius yra kompetentingas ir patikimas, jie linkę labiau pasitikėti pačiu turiniu. Autorystės aiškumas, pavyzdžiui, matant autoriaus vardą, pareigas ir patirtį, suteikia galimybę įvertinti jo ekspertiškumą ir patikimumą. Jei šios savybės suvokiamos teigiamai, didėja ir turinio patikimumas vartotojų akyse (Jenkins et al., 2020b). Į anonimiškai pateiktą turinį vartotojai žvelgia skeptiškiau, nei į tą, kurio autorystė yra atskleista. Žinutės skaitmeninėje erdvėje, kurios yra grindžiamos anoniminiais šaltiniais, dažnai laikomos mažiau patikimomis nei tos, kuriose šaltiniai įvardyti (Smith & Matthews, 2021). Skaidrumo stoka kelia natūralų nepasitikėjimą – skaitytojai nežino, kas pateikia informaciją, todėl abejoja jos autoritetumu. Kita vertus, ne visada anonimiškumas visiškai sugriaua pasitikėjimą. Pavyzdžiui, Schlosser (2020) atskleidė, kad anoniminiai šaltiniai tam tikromis aplinkybėmis gali būti vertinami taip pat patikimai kaip ir tie, kurių autorystė buvo atskleista. Tai rodo, jog kontekstas ir anonimiškumo priežastys svarbios – jei auditorija suvokia anonimiškumą kaip pateisinamą (pavyzdžiui, norint atvirai pasidalinti jautria patirtimi), pasitikėjimas labai nenukentės. Vis dėlto bendroje socialinių tinklų aplinkoje galioja tendencija: didesnis autorystės aiškumas paprastai didina pasitikėjimą, o visiškas anonimiškumas jį mažina, ypač jei nėra kitų patikimumo užuominų.

Svarbu ir tai, kaip atskleidžiama autorystė. Socialiniai tinklai taiko įvairius autorystės patvirtinimo būdus – nuo paprasto vardo paviešinimo iki patvirtintos paskyros ženklelio. Patvirtinimo ženklai, pavyzdžiui, „mėlynas varnelės“ statusas prie profilio, gali stipriai paveikti vartotojų pasitikėjimą. Geels'o et al. (2024) eksperimentas parodė, kad vartotojai dažnai sutapatina šaltinio autentifikavimą su žinutės patikimumu – jei profilis pažymėtas kaip patvirtintas, manoma, jog ir jo skleidžiama informacija patikima. Šis efektas vadinamas „virtualaus balto chalato“ (angl. *virtual lab coat*) efektu: žinutės iš paskyrų su patvirtintu statusu buvo laikomos patikimesnėmis net ir tuomet, kai tas statusas (pavyzdžiui, profesinė kvalifikacija) visiškai nesusijęs su žinutės tema. Autorystės aiškumas (pavyzdžiui, matomas vardas ar pažymėjimas) veikia kaip heuristika – lengvas kelias auditorijai spręsti apie turinio patikimumą. Nors tai padidina pasitikėjimą, kyla rizika, kad vartotojai nebevertins turinio kritiškai ir gali apsigauti, jei pats autorystės indikatorius klaidinantis. Pavyzdžiui, socialiniame tinkle „Twitter“ (dabar „X“) mėlynas ženklelis prie paskyros pavadinimo reiškė, kad platforma patvirtino paskyros, t. y. žinomų asmenų ar prekės ženklų, autentiškumą. Vartotojai ilginiui priprato manyti, kad jei ženklelis yra – šaltinis patikimas. 2022 m. pabaigoje „Twitter“ pakeitus politiką ir leidus įsigyti patvirtinimo ženklelį už mokestį be tapatybės tikrinimo, pasipylė apsimetinėjimo incidentai. Vos per kelias dienas netikri, apsimetę įmonėmis profiliai, turintys nupirktą ženklelį, prirašė melagingų žinučių, kurias daugelis vartotojų palaikė tikromis, manydami, jog šaltinis oficialus. Kai patikimumo indikatorius, šiuo atveju ženklelis, praranda prasmę, vartotojai suklaidinami, o pasekmės daro realią įtaką įmonių finansiniams rezultatams – kai kurių įmonių akcijų

kainos krito dėl melagingų pranešimų (Geels et al., 2024). Skaidri ir patikima autentiškumo patvirtinimo sistema yra kritiškai svarbi siekiant išlaikyti pasitikėjimą turiniu socialiniuose tinkluose.

Diskusijų platformose, tokiose kaip „Reddit“, ar anoniminiuose forumuose, dažnai leidžiama bendrauti pseudonimiškai. Nors tai skatina atvirumą ir leidžia kalbėti jautriomis temomis be baimės dėl asmeninės tapatybės atskleidimo, bendras pasitikėjimo lygis tokiuose tinkluose gali būti mažesnis. Tyrimai rodo, kad anoniminėje aplinkoje išauga nekonstruktyvių ar agresyvių pasisakymų tikimybė, o tai netiesiogiai kenkia vartotojų pasitikėjimui – jei bendravimas priešiškas ir pilnas dezinformacijos, vartotojai sunkiau pasitiki ten randama informacija (Sardá et al., 2019). Atkreipdamos į tai dėmesį, platformos, anksčiau leidusios visišką anonimiškumą, svarsto kompromisinius sprendimus. Pavyzdžiui, siūlomos sistemos, kur vartotojo tapatybė patikrinama užkulisiuose, tačiau viešai skelbiamas pseudonimas (Gerrard, 2021). Tokiu būdu siekiama derinti anonimiškumo naudą (privatumą, saviraišką) su tam tikru atsakomybės mechanizmu, kad kilus piktnaudžiavimui būtų įmanoma nustatyti autorių. Vartotojai palankiai vertina anonimiškumą tik tuomet, kai mano, kad pats viešinamas turinys yra patikimas. Taip vėl grįžtama prie esminio pastebėjimo: pasitikėjimą lemia turinio kokybė ir šaltinio patikimumo suvokimas, ne vien techninis vardo parodymas ar paslėpimas.

Skaidrumas skatina pasitikėjimą, nes sukuria atskaitomybės jausmą. Kai autorius nepaslėptas, tikimasi, jog jis atsako už savo žodžius – tai psichologiškai ramina auditoriją. Priešingai, anonimiškumas gali mažinti pasitikėjimą ne tik konkrečiu turiniu, bet ir pačia platforma, jei joje daug neaiškios kilmės informacijos. Dėl to socialinių tinklų naudotojai vertina aiškią autorystę: auditorija daugiau pasitiki turiniu, kuris siejamas su aiškiai atpažįstamu ir reputaciją turinčiu šaltiniu. Net ir paprastas autoriaus „autoritetas“ (statusas) turi reikšmę – oficialių ar autoritetingų šaltinių buvimas ženkliai didina šaltinio patikimumo vertinimus (Yaqub et al., 2020).

Šiuolaikinė praktika skaitmeninėje erdvėje kelia naują autorystės iššūkį – turinį, generuojamą dirbtinio intelekto. DI turinio autorystės dilema kyla dėl to, kad nėra aiškaus žmogiško autoriaus: vartotojai gali klausti, ar galima pasitikėti žinute, jei ją parašė ne žmogus? Viena vertus, skaidrumas reikalautų atskleisti, kad turinys generuotas DI, tačiau Buchanan ir Hickman'as (2024) pastebi, jog vien toks atskleidimas gali sumažinti pasitikėjimą turiniu. Pavyzdžiui, Kanzaso universiteto eksperimentas, kuriame naujienų straipsnių skaitytojams buvo nurodoma, kiek prisidėjo DI, atskleidė, kad skaitytojai menkliau pasitiki turiniu, jei mano, kad jį rengiant kažkoku būdu dalyvavo DI (The University of Kansas, 2024). Net ir nevisiškai suprasdami DI indėlį, žmonės jautėsi mažiau pasitikintys straipsnio turiniu vien dėl DI paminėjimo – tai atskleidžia visuomenės nusistatymą, kad žmogaus parašytas turinys patikimesnis nei dirbtinio intelekto. Turbūt taip yra todėl, nes žmogus yra atsakingas už savo darbą, o DI, savo ruožtu, yra naujovė, kuri nesupranta, kada padaro klaidas. Kita vertus, Li ir Yang (2024) eksperimentas parodė, jog turinio ženklinimas kaip „sugeneruoto DI“ pats savaime nebūtinai mažina pasitikėjimą. Buvo nustatyta, kad aiškus DI turinio pažymėjimas neturėjo reikšmingos įtakos vartotojų suvokiamam žinutės tikslumui ar patikimumui. Žmonių polinkis tikėti ar dalintis žinute nebuvo statistiškai sumažintas dėl dirbtinio intelekto sugeneruoto turinio (angl. *AI-Generated Content*) etiketės. Ši dilema verčia spręsti, kaip geriausia elgtis socialinių tinklų platformoms: ar žymėti DI kurtą turinį, taip užtikrinant skaidrumą, bet rizikuojant neigiamu turinio vertinimu, ar nežymint ir taip galimai klaidinant dėl autorystės.

Apibendrinant, tiek teoriniai modeliai, tiek empiriniai tyrimai patvirtina, jog turinio autoriaus identifikavimas yra esminis veiksnys vartotojų pasitikėjimui ir įsitraukimui skaitmeninėje erdvėje.

Skaidri autorystė paprastai stiprina pasitikėjimą ir palankumą, nes leidžia auditorijai pasitelkti socialinius signalus – ekspertiškumą, reputaciją, autentiškumą – vertinant turinį. Tuo tarpu anonimiškumas ar neaiški autorystė dažnai sietina su skepticizmu, nebent egzistuoja kiti pasitikėjimą kuriantys mechanizmai. Dabartinėje skaitmeninėje ekosistemoje, kur atsiranda ir DI generuotas turinys, autorystės klausimas darosi dar aktualesnis: formuojasi naujos normos ir politikos, siekiančios suderinti skaidrumo principą su vartotojų poreikiu turėti patikimus turinio indikatorius.

2.4. Vartotojų pasitikėjimas bei ketinimai dirbtinio intelekto ir žmogaus kuriamo turinio atveju

Vertindami dirbtinio intelekto sukurtą turinį, vartotojai dažnai lygina jį su žmogaus kurtu turiniu. Tarp šių dviejų turinio tipų pastebimi skirtumai patikimumo (angl. *credibility*), autentiškumo (angl. *authenticity*) pojūčio bei vartotojų įsitraukimo (angl. *engagement*) lygmenyse. Konkretus turinio vertinimas priklauso nuo jo rūšies ir konteksto – pavyzdžiui, ar tai naujienų straipsnis, ar pramoginis įrašas, ar reklamų generuojamas tekstas (Joy et al., 2024). Daugelis vartotojų iš prigimties yra šiek tiek skeptiški DI turinio atžvilgiu. Vien etiketė, imponuojanti, kad turinys buvo sukurtas naudojant dirbtinį intelektą, gali sumažinti turinio patikimumą vartotojo požiūriu (Altay & Gilardi, 2024a). Kita vertus, kai kuriais atvejais vartotojai pripažįsta, kad DI generuotas turinys gali būti objektyvus ir tikslus – ypač kai kalbama apie faktinę, analitinę informaciją. Pavyzdžiui, PennState (2022) tyrimas moderavimo srityje parodė, jog vartotojai gali pasitikėti DI taip pat, kaip ir žmonėmis, jei tiki DI sistemos tikslumu ir objektyvumu. Patikimumo suvokimas priklauso nuo vartotojo nuostatų – tie, kurie mano, kad DI yra tikslus, linkę juo pasitikėti labiau. Pasitikėjimui DI turiniu didelę įtaką daro ir paties šaltinio autoritetingumas. Hofeditz'o et al. (2021) tyrimas atskleidė, kad turinio patikimumą labiausiai lemia žiniasklaidos prekės ženklo patikimumas ir vartotojų patirtis, o ne tai, ar tekstas parašytas DI. Pasitikėjimo skirtumai DI ir žmogaus kuriamu turiniu gali būti mažiau akivaizdūs, kai atkreipiamas dėmesys į tai, kas skelbia turinį, kokiame kontekste jis yra skelbiamas ir koks paties vartotojo technologijų priėmimo lygis. Vis dėlto, kol kas vyrauja tendencija manyti, kad žmogaus kurtas turinys patikimesnis, tačiau augant DI galimybėms ir vartotojams priimant jį, atotrūkis gali mažėti – ypač jei DI turinys kuriamas kokybiškai, su žmogaus priežiūra.

Turinio autentiškumo – įspūdžio, kad turinys yra tikras, nuoširdus, atspindintis realią žmogiškąją patirtį – požiūriu, žmogaus kuriamas turinys neabejotinai turi pranašumą vartotojų akyse. Net 59 proc. vartotojų autentiškiausiu laiko kitų vartotojų sukurtą turinį (angl. *user-generated content*), tuo tarpu tik 19 proc. mano, kad autentiškiausias yra prekės ženklo, t. y. konkrečių įmonių, kurtas turinys (Nosto, 2021). DI sukurtas turinys kol kas neretai prilyginamas negyvam, korporaciniam turiniui, stokojančiam žmogaus emocijos. Vartotojai gali jausti, kad DI tekste ar vaizde trūksta žmogiškosios dvasios (angl. *human essence*) – unikalios perspektyvos, jausmo, kurį suteikia žmogaus kūryba (Kim et al., 2015). Vis dėlto, DI turinys gali atrodyti autentiškas, jei jis imituoja realius žmones ir jų kūrybą. DI grindžiamos vaizdo klastotės (angl. *deepfake*) technologijos DI pagalba gali sugeneruoti visiškai fotorealistiškus vaizdus ir vaizdo įrašus. 2023 m. socialiniuose tinkluose plačiai paplito popiežiaus Pranciškaus su balta pūkine striuke nuotrauka. Tūkstančiai žmonių patikėjo šios nuotraukos tikrumu, kol paaiškėjo, kad tai – DI sugeneruotas vaizdas, sukurtas vaizdo generatoriaus „Midjourney“ pagalba (žr. 2 pav.).



2 pav. „Midjourney“ sukurta vaizdo klastotė (autorius nežinomas)

Žurnalas „Time“ šį įvykį praminė „pirmuoju tikrai virusiniu dezinformacijos įvykiu, paskatintu vaizdo klastojimo technologijos“ (Campos et al., 2023). Šis DI sukurta vaizdas buvo toks tikroviškas, kad apgavo itin daug žmonių. Analizuojant šį konkretų atvejį, galima išskirti du svarbius pastebėjimus: pirmą – DI gali sukurti itin realiai atrodančius fiktyvius vaizdus, antra – žmonės, sužinoję apgaulę, pasijuto suklaidinti ir dėl to krito jų pasitikėjimas turiniu pasidalinusiomis platformomis. Skaitmeninėje erdvėje kilo diskusijos dėl tokio ir panašaus turinio žymėjimo ir tikrinimo, vartotojai suprato, kad akimis neužtenka pasikliauti, norint nustatyti turinio autentiškumą (Campos et al., 2023). Dėl šios priežasties, žmogaus kuriamas turinys laikomas autentišku, tuo tarpu DI kuriamas turinys asocijuojamas su sintetiškumu. Vis dėlto, jei žinoma turinio kilmė, t. y. kas yra jo autorius, ir ar sukurta turinys atliepia vartotojo poreikius, jis gali būti priimamas. *Deloitte* (2024) tyrimas atskleidė, kad nemaža dalis socialinių tinklų vartotojų mano, jog DI sukurta turinys jau gali prilygti arba net pranokti žmogaus kurtą turinį pagal personalizavimo galimybes ir pramoginę vertę. Ši tendencija itin ryški tarp jaunesnių, su technologijomis labiau susijusių kartų.

Įsitraukimas socialiniuose tinkluose apibrėžiamas kaip vartotojų veiksmai sąveikaujant su turiniu: peržiūros, paspaudimai, komentarai, pasidalijimai, reakcijos. Svarbu atkreipti dėmesį, kad vartotojų elgesys ne visada tiesiogiai atspindi pasitikėjimą (Lim & Rasul, 2022). Pavyzdžiui, sensacingas ar pramoginis DI turinys gali sulaukti milžiniško įsitraukimo net jei juo nevysiškai pasitikima – vartotojai gali dalintis juo dėl pramogos ar šoko efekto. Vartotojai linkę dalintis straipsniais, kurie, jų manymu, yra patikimi ir vertingi. Jei turinys atrodo parašytas eksperto ar liudija autentišką patirtį, komentarai ir pasidalijimai būna pozityvesni, diskusijos gilesnės. DI generuotas turinys, ypač, jei tai be konteksto pateikti faktai ar automatinės naujienos, gali sulaukti mažesnio aktyvaus įsitraukimo. Pavyzdžiui, Altay‘aus ir Gilardi‘o (2024) tyrimas nustatė, kad pažymėjus naujienas kaip generuotas DI, vartotojų ketinimas jomis dalytis sumažėjo. Žmonės nenorėjo platinti turinio, kuriuo suabejojo dėl DI žymos, net jei tas turinys galėjo būti teisingas. Šiuo atveju pasitikėjimas tiesiogiai veikia dalijimosi ketinimus – jei vartotojas turiniu nepasitiki, jis jo neplatins. Kita vertus, jei DI turinys yra itin įtraukiantis arba intriguojantis, vartotojai gali aktyviai su juo sąveikauti net ir suvokdami jo kilmę. Grįžtant prie popiežiaus pavyzdžio, žmonės itin aktyviai juo dalijosi iki tol, kol tikroji atvaizdo kilmė

nebuvo atskleista. Klaidinantis DI turinys gali paskatinti didžiulį įsitraukimą apgaulės būdu. Dėl šios priežasties socialinių tinklų platformos turėtų atkreipti dėmesį į tai, kaip valdyti DI turinio sklaidą, kad vartotojų įsitraukimas nevestų prie nekontroliuojamo dezinformacijos plitimo (Campos et al., 2023). Kalbant apie vartotojų įsitraukimą su patikimu turiniu, verta paminėti, kad DI gali pasitarnauti teigiamai, kai naudojamas turiniui personalizuoti. Asmeniškai pritaikytas, DI atrinktas turinys, pavyzdžiui, muzikos rekomendacijos „Spotify“ platformoje, vazido įrašai „YouTube“, gali didinti vartotojų įsitraukimą, jei vartotojas pasitiki rekomendacijų sistema. Jei DI sistema gerai „pažįsta“ vartotojo pomėgius ir siūlo jam aktualų žmogaus ar DI kurtą turinį, vartotojas daugiau laiko praleis platformoje, labiau įsitrauks. Čia pasitikėjimas veikia tarp eilučių – vartotojas gal nebūtinai galvoja, „ar galiu pasitikėti šia rekomendacija?“, bet jei kelis kartus gavo prastą ar neaktualių pasiūlymų, jo pasitikėjimas sistema mažės ir įsitraukimas kris (Piduru, 2023). Tad skirtumas tarp DI ir žmogaus turinio įsitraukimo gali pasireikšti per kokybę: kokybiškas turinys – nesvarbu, sukurtas žmogaus ar DI, – įtrauks, o nekokybiškas (pavyzdžiui, paviršutiniškas DI sugeneruotas tekstas) liks vartotojų ignoruojamas. Tam tikrose srityse vartotojai vertina DI turinį kaip papildantį žmogiškąjį, o nebūtinai konkuruojantį, pavyzdžiui, klientų aptarnavimo pokalbiuose socialiniuose tinkluose kai kurie žmonės neretai sutinka bendrauti su pokalbių robotu (angl. *chatbot*) ir gauna greitus atsakymus – tai didina jų pasitenkinimą ir įsitraukimą, net jei žino, kad bendrauja su DI, tačiau svarbu, kad DI atsakymai būtų tikslūs (Lappeman et al., 2023). Žmogaus kuriamas turinys vis dar laikomas standartu autentiškumui ir paprastai sulaukia didesnio pasitikėjimu grįsto įsitraukimo, tačiau DI turinys sparčiai tobulėja, o vartotojų įsitraukimas, laikui bėgant, labiau priklausys nuo turinio kokybės bei konteksto, nei vien nuo to, kas jį sukūrė. Kai DI turinys yra akivaizdžiai naudingas, pavyzdžiui, automatiškai sugeneruotos asmeninės suvestinės, kurios sutaupo laiko, vartotojai noriai jį naudoja – čia įsitraukimą lemia suvokiama nauda ir patogumas.

Toliau pateikiamoje lentelėje (žr. 11 lentelę) susisteminami pagrindiniai panašumai ir skirtumai kalbant apie DI ir žmogaus kuriamą turinį socialiniuose tinkluose bei vartotojų elgseną. Išskiriama bendra sritis, pastebėjimai bei keletas literatūros pavyzdžių.

11 lentelė. Vartotojų elgsena DI ir žmogaus kuriamo turinio atveju

Aspektas	Literatūros pavyzdžiai
Patikintukų paspaudimas (angl. <i>likes</i>)	(Altay & Gilardi, 2024; Huschens et al., 2023)
DI: vartotojai dažnai vertina DI parašytą tekstą kaip itin aiškų ir įtraukiantį – kartais net labiau nei žmogaus rašytą turinį, o tai rodo, kad jis gali pelnyti žmonių pripažinimą. Tačiau kai turinys aiškiai pažymimas kaip sukurtas su DI pagalba, vartotojai gali vengti paspausti „patinka“, greičiausiai dėl vyraujančio skepticizmo.	
Žmogus: kadangi nėra jokios stigmos prieš žmogaus kuriamą turinį, kokybiškas turinys greičiausiai sulauks palaikymo, t. y. paspaudimo „patinka“ dėl savo vertės, be abejonių, kurias gali kelti dirbtinio intelekto sukurtas turinys.	
Dalijimasis (angl. <i>sharing</i>)	(Altay & Gilardi, 2024; M. Yang ir Rau, 2024)
DI: vartotojų noras dalintis dirbtinio intelekto sukurtu turiniu gali sumažėti, jei atskleidžiama jo kilmė – nepriklausomai nuo antraštės tikrumo ar autoriaus. Vis dėlto, kai kurie autoriai tikina, kad autorystės atskleidimas realioje praktikoje kartais daro tik minimalų poveikį dalijimosi ketinimams. Tai leidžia manyti, kad vartotojai gali dalintis DI turiniu panašiai kaip ir žmogaus sukurtu turiniu, jei kiti veiksniai, pavyzdžiui, kokybė ar aktualumas, yra pakankamai svarbūs.	
Žmogus: turinys, kuris yra žinomas kaip žmogaus parašytas, nesusiduria su skepticizmo barjeru, su kuriuo gali susidurti AI generuotas turinys, todėl didžiausią įtaką jo dalijimosi potencialui daro kokybė ir aktualumas.	
Peržiūros ir paspaudimai (angl. <i>viewing and click-through</i>)	(Graefe & Bohlken, 2020)

Aspektas	Literatūros pavyzdžiai
DI: žmonės DI turinį peržiūrės, jei jis atrodys naudingas ar įdomus, tačiau kiti gali vengti paspausti ant DI sukurtų nuorodų dėl pasitikėjimo stokos. Kita vertus, jei DI sukurtas turinys yra aukštos kokybės ir nėra aiškiai pažymėtas, vartotojai gali su juo sąveikauti taip pat kaip ir su žmogaus kurtu turiniu.	
Žmogus: žmonių sukurtas turinys dažnai atitinka vartotojų lūkesčius ir paprastai sulaukia išankstinio pasitikėjimo. Vartotojai linkę spausti ir peržiūrėti žmogaus parašytus straipsnius, ypač jei tema juos domina.	
Komentariai, skundai, išsaugojimai (angl. <i>comments, reports, saves</i>)	(Niu Yongbo, 2025)
DI: šiuo metu akademinių duomenų apie tai, kaip vartotojai komentuoja, praneša apie pažeidimus ar išsaugo DI sukurtus įrašus, palyginti su žmonių sukurtais, yra nedaug. Daroma prielaida, kad jei vartotojai atpažįsta turinį kaip sugeneruotą DI pagalba, jie gali būti labiau linkę skeptiškai ar neigiamai komentuoti bei greičiau pranešti apie klaidinančią informaciją. Taip pat tikėtina, kad jie rečiau „išsaugos“ ar pažymės DI turinį kaip vertą išsaugojimo, jei kils abejonių dėl jo patikimumo. Tačiau šie vartotojų elgsenos aspektai dar nėra išsamiai ištirti ir dabartinė mokslinė literatūra šioje srityje neturi galutinių išvadų.	
Žmogus: patikimas ar įtraukiantis žmogaus parašytas įrašas sulauks komentarų ir išsaugojimų, tuo tarpu žalingas ar prastos kokybės turinys gali būti praneštas kaip netinkamas – šie veiksmai priklauso nuo pačio turinio esmės.	

Toliau pateikiami (žr. 12 lentelę) pagrindiniai pastebėjimai kalbant apie pasitikėjimą DI ir jo kuriamu turiniu, lyginant juos iš patikimumo, autentiškumo, emocinės įtraukties ir skaidrumo perspektyvų. Išskiriama sritis, keletas literatūros pavyzdžių bei pagrindiniai pastebėjimai.

12 lentelė. Vartotojų pasitikėjimas DI ir žmogaus kuriamu turiniu

Aspektas	Literatūros pavyzdžiai
Suvokiamas patikimumas (angl. <i>perceived credibility</i>)	(Graefe & Bohlken, 2020; Huschens et al., 2023; Kim et al., 2015)
DI: apskritai vartotojai gali laikyti DI sukurtą informaciją tokia pat patikima kaip ir žmonių kurtą, jei turinys yra aukštos kokybės ir DI kilmė nėra aiškiai pabrėžta. Kita vertus, viešasis skepticizmas vis dar egzistuoja – pavyzdžiui, AI parašyti naujienų straipsniai tam tikruose kontekstuose buvo vertinami kaip mažiau patikimi nei analogiški žmonių sukurti straipsniai. Šiuose suvokimuose svarbų vaidmenį atlieka skaidrumas.	
Žmogus: vartotojai dažnai iš anksto daro prielaidą, kad žmonių parašytos naujienos ar įrašai yra patikimi. Jei vartotojai tiki, jog turinys buvo sukurtas žmogaus, net jei iš tikrųjų taip nebuvo, jie dažnai jį vertina kaip patikimesnį ir aukštesnės kokybės.	
Suvokiamas autentiškumas	(Eigenraam et al., 2021; Kim et al., 2015; J. Yang et al., 2021)
DI: DI sukurtas turinys dažnai laikomas stokojančiu „žmogiškojo faktoriaus“. Vartotojai yra linkę jį vertinti kaip mažiau autentišką ar asmenišką, ypač jei žino, kad jį sukūrė DI sistema.	
Žmogus: žmonių sukurtas turinys paprastai laikomas autentišku, nes už jo slypi realus asmuo. Žmonių kūrėjai gali perteikti asmeninę patirtį, emocijas ir spontaniškumą, todėl jų turinys atrodo tikresnis.	
Emocinė įtrauktis (angl. <i>emotional engagement</i>)	(M. Yang & Rau, 2024; Niu Yongbo, 2025; Pusztahelyi, s.a.)
DI: dirbtinio intelekto sukurtam turiniui sunkiai sekasi sukelti tokias pačias emocines reakcijas kaip žmonių kūrybai. Nors DI generuojamas tekstas ir vaizdai gali būti informatyvūs ar pramoginiai, jie dažnai stokoja gilaus emocinio išraiškingumo, empatijos ar asmenybės, kurie stipriai rezonuoja su auditorija.	
Žmogus: žmonių sukurtas turinys paprastai sukelia stipresnį emocinį ryšį, nes žmonės geba perteikti jausmus, humorą ir artumą. Žurnalistai, atlikėjai ar menininkai gali pasitelkti asmeninę patirtį ir emocinį intelektą, kad užmegztų glaudesnę santykį su auditorija.	
Skaidrumas (angl. <i>transparency</i>)	(Altay & Gilardi, 2024a; M. Yang & Rau, 2024)
DI: atskleidimas, jog turinį sukūrė DI gali turėti dvilypį poveikį pasitikėjimui. Iš vienos pusės, skaidrumas padeda išvengti vartotojų apgaulės jausmo ir užtikrina sąžiningumą, tačiau, kita vertus, toks atskleidimas dažnai sukelia skeptiškesnę reakciją – vartotojai tampa sąmoningesni dėl galimų DI algoritmų apribojimų ar šališkumo, dėl ko gali sumažėti jų pasitikėjimas ar įsitraukimas.	

Aspektas	Literatūros pavyzdžiai
Žmogus: žmonių sukurtam turiniui skaidrumas dažniausiai nėra problema, nes autoriaus žmogiškoji kilmė laikoma numatyta prielaida. Paprastai nereikia aiškios žymos, kad turinys buvo „parašytas žmogaus“. Skirtingai nei DI turinio autorystės atskleidimas, „žmogiškos kilmės“ pažyma retai mažina vartotojų įsitraukimą – priešingai, ji gali padidinti patikimumą, nes patvirtina autoriaus atsakomybę už pateikiamą informaciją.	

Vartotojų pasitikėjimas – tai kertinis socialinių tinklų ekosistemos aspektas, palaikantis kokybiškos informacijos sklaidą. Vartotojų pasitikėjimas susijęs su tokiais veiksniais kaip skaidrumas, sąžiningumas, privatumas ir autorystės atskleidimas. DI generuojamas turinys skatina permąstyti tradicinius pasitikėjimo pagrindus: vartotojai nori žinoti, kas sukūrė turinį, kaip jis kurtas, ar procesas buvo skaidrus ir teisingas. Jei šie lūkesčiai tenkinami, DI turinys gali būti priimtas; jei ne – kyla skepticizmas.

Apibendrinant, skirtumai tarp DI ir žmogaus kurto turinio vertinimo dar gana ryškūs. Žmonių kurtas turinys yra laikomas patikimesniu ir autentiškesniu, vartotojai linkę labiau juo pasitikėti ir su juo sąveikauti. DI turinys dažnai sulaukia padidinto dėmesio, bet taip pat ir padidinto įtarumo – vartotojai gali žavėtis DI, tačiau kartu abejoti turinio patikimumu. Vartotojų elgsena DI kuriamo turinio kontekste svyruoja nuo perdėto patiklumo iki perdėto skepticizmo kaip ženklavimo eksperimente. Tai rodo, jog visuomenė dar mokosi prisitaikyti prie naujos realybės, kurioje turinio autorius gali būti nebe žmogus. Ateityje galime tikėtis, kad vartotojų pasitikėjimo dinamika evoliucionuos. Gerėjant DI turinio kokybei ir diegiant atsakingo naudojimo priemones, vartotojai, tikėtina, ims vertinti turinį ne vien pagal tai, kas jį paruošė, bet ir pagal jo kokybę. Vis dėlto, norint to pasiekti, reikės stiprių pastangų užtikrinant etišką DI naudojimą. Etiniai ir teisiniai sprendimai čia atlieka lemiamą vaidmenį: aiškos taisyklės dėl DI turinio žymėjimo, atsakomybės ir skaidrumo yra būtinos, siekiant užtikrinti, kad DI neklaidintų, o priešingai, padėtų vartotojams.

2.5. Dirbtinio intelekto ir žmogaus kuriamo turinio poveikio vartotojų pasitikėjimui ir ketinimams konceptualaus modelio pagrindimas

Atsižvelgiant į ankstesnėse dalyse aptartą teorinę literatūrą, šiame skyriuje pateikiamas konceptualus dirbtinio intelekto ir žmogaus kuriamo turinio poveikio vartotojų pasitikėjimui ir ketinimams modelis (žr. 3 pav.), iliustruojantis, kaip dirbtinio intelekto ir žmogaus kuriamas turinys socialiniuose tinkluose gali paveikti vartotojų pasitikėjimą ir jų ketinimus.

Autorių vaidmuo turinio vertinime ir pasitikėjime. Remiantis ankstesniuose skyriuose aptartomis įžvalgomis, dirbtinio intelekto ir žmogaus kuriamo turinio poveikį vartotojų pasitikėjimui bei sąveikos ketinimams lemia keli pagrindiniai komponentai ir jų tarpusavio ryšiai. Literatūroje pabrėžiama, kad pats turinio autorius – ar tai DI, ar žmogus – stipriai veikia vartotojų suvokimą apie turinio autentiškumą, kokybę ir patikimumą (Berber Sardinha, 2024; Buchanan & Hickman, 2024; Graefe & Bohlken, 2020; Huschens et al., 2023; Ismagilova et al., 2020; Li & Yang, 2024; Mirbabaie & Stieglitz, 2021). Šie suvokimai tiesiogiai siejami su vartotojų pasitikėjimo lygiu ir ketinimu sąveikauti su turiniu (Belanche et al., 2024). Kitaip tariant, autorystė tampa esminiu veiksnio: turinio kilmė – parašyta žmogaus ar sugeneruota DI – nulemia pradinius vartotojo lūkesčius ir vertinimą, kurie atsispindi pasitikėjime bei ketinimuose.

Autorystės suvokimas ir turinio kilmė. Žmogaus sukurtas turinys socialiniuose tinkluose dažnai siejamas su asmeniškumu ir „tikrumu“, todėl vartotojai juo linkę pasitikėti labiau nei DI sugeneruotu turiniu (Huschens et al., 2023). DI generuojamas turinys vertinamas skeptiškiau dėl jo mechaninio,

sintetinio pobūdžio (X. Chen et al., 2022; Eigenraam et al., 2021; Lim & Rasul, 2022). Pavyzdžiui, DI turinys gali būti objektyvesnis ir faktais grįstas, tačiau stokoja emocinio ryšio – vartotojai pastebi, kad virtualūs nuomonės formuotojai pateikia naudingas faktines rekomendacijas, bet labiau įsitraukia į žmogaus kurtą turinį dėl emocinio atspalvio (Belanche et al., 2024). Vis dėlto, kai turinio autorystė nėra aiškiai atskleista, vartotojai priversti patys suvokti ar spėti, kas yra turinio kūrėjas. Jie pasitelkia tam tikras užuominas – teksto toną, raiškos stilių, vizualinę kokybę – periferinius signalus, leidžiančius numanyti turinio šaltinį (Jenkins et al., 2020). Toks autorystės suvokimas gali skirtis nuo realybės ir paveikti tolimesnį vertinimą: jei turinys atrodo nenatūralus ar neatitinka tipinių žmogaus komunikacijos bruožų, vartotojams kyla įtarimų dėl jo kilmės. Tyrimai rodo, kad net menkiausi turinio ir jo pateikimo neatitikimai vartotojams signalizuoja neautentiškumą bei manipuliatyvumą, dėl ko jie praranda pasitikėjimą (Huschens et al., 2023; Ismagilova et al., 2020; Li & Yang, 2024; Mirbabaie & Stieglitz, 2021). Ypač DI kuriami vaizdai ar tekstai, turintys „sintetinių“ požymių, gali sukelti keistumo jausmą (angl. *uncanny valley effect*) ir mažinti pasitikėjimą (Belanche et al., 2024; Gu et al., 2024; Sargin, 2024). Verta paminėti, kad tobulėjant DI modeliams, DI generuojamas turinys darosi vis panašesnis į žmogaus kurtą turinį, todėl vartotojams tampa vis sunkiau atpažinti autorystę ir užtikrinti skaidrumą (Park et al., 2024). Taigi, vartotojų autorystės suvokimas – tai, ką jie mano apie turinio kilmę – yra kritinis tarpinis veiksnys, galintis sustiprinti arba susilpninti turinio poveikį: jei vartotojas įtaria, kad turinį sukūrė DI, jis gali iš anksto žiūrėti skeptiškiau, o jei mano, kad autorius žmogus, turinys gali atrodyti patikimesnis.

Autorystės atskleidimas, skaidrumas. Autorystės atskleidimas – tai aiškus informacijos suteikimas apie tai, kas yra turinio kūrėjas nurodant autorių, ar pranešant, kad turinį generavo DI. Skaidrus autorystės nurodymas literatūroje siejamas su didesniu vartotojų pasitikėjimu (Larsson & Heintz, 2020; Wang, 2022). Pagal šaltinio patikimumo modelį (angl. *Source Credibility Model*), auditorijos pasitikėjimas turiniu labai priklauso nuo pačio autoriaus patikimumo vertinimo (Seiler & Kucza, 2017). Kai vartotojai žino turinio autorių ir mano jį esant kompetentingą bei patikimą, jie labiau linkę tikėti ir pačiu turiniu (Geels et al., 2024; Kim et al., 2015). Aiškiai matoma autorystė, pvz., atskleistas autoriaus vardas, pareigos, suteikia kontekstą – galima įvertinti šaltinio reputaciją, o tai stiprina pasitikėjimą. Priešingai, anonimiškas ar neaiškios kilmės turinys kelia daugiau abejonių: nepažįstamas, neatskleistas šaltinis vartotojams sunkiau įvertinamas, todėl toks turinys dažnai laikomas mažiau patikimu. Empiriniai tyrimai patvirtina, kad vartotojai žvelgia skeptiškiau į žinutes, kurių autorystė neatskleista, lyginant su tomis, kuriose šaltinis įvardytas. Skaidrumo trūkumas iššaukia natūralų nepasitikėjimą – jei skaitytojas nežino, kas pateikė informaciją, jis abejoja jos autoritetingu (Ferrario et al., 2020; Geels et al., 2024). Vis dėlto literatūroje taip pat pažymima, kad autorystės atskleidimo poveikis nėra vienareikšmis visais atvejais: tam tikromis aplinkybėmis anonimiškumas gali nesumažinti pasitikėjimo, pvz., jei auditorija suvokia anonimiškumą kaip pateisinamą dėl dalijimosi itin jautria patirtimi (Smith & Matthews, 2021). Nepaisant tokių išimčių, bendra tendencija socialinių tinklų komunikacijoje yra tokia, kad didesnis skaidrumas, pavyzdžiui, aiškus autorystės atskleidimas, paprastai daro teigiamą įtaką pasitikėjimui (Felzmann et al., 2020; Singhal et al., 2024; Wang, 2022). Dėl šios priežasties konceptualiaame modelyje autorystės atskleidimas laikomas svarbiu moderuojančiu veiksniu – jis gali sustiprinti arba susilpninti kitų komponentų poveikį pasitikėjimui.

Turinio vertinimas per jo autentiškumą ir kokybę. Nepriklausomai nuo to, kas kuria turinį, vartotojai jį įvertina pagal tam tikrus kriterijus, iš kurių svarbiausi šio tyrimo kontekste yra suvokiamas turinio autentiškumas ir kokybė. Turinio autentiškumo ir kokybės suvokimas literatūroje

išskiriamas kaip esminis elementas vartotojų pasitikėjimo formavimuisi. Vartotojai linkę labiau pasitikėti turiniu, kuris jiems atrodo tikras ir kokybiškas (Eigenraam et al., 2021; Kim et al., 2015; J. Yang et al., 2021). Autentiškas turinys šio tyrimo kontekste apibrėžiamas kaip pasižymintis sąžiningumu, nuoseklumu, gebėjimu užmegzti emocinį ryšį, ir neturintis „nenatūralių“ – sintetiskų ar keistų – bruožų (Jenkins et al., 2020; Kemp et al., 2021; Lee, 2020; Maares et al., 2021; Schreiner et al., 2021). Jei vartotojas pastebi turinyje neatitikimų, nenatūralumo ar dirbtinumo, kyla įtarimų dėl turinio tikslumo ir patikimumo, ir pasitikėjimas gali būti prarastas (Gu et al., 2024). Pavyzdžiui, tyrimai parodė, kad turinio nuoseklumas ir sąžiningumas vartotojų akyse yra pagrindiniai autentiškumo rodikliai – jų nebuvimas ar klaidinantys elementai sumažina pasitikėjimą turiniu (Huschens et al., 2023; Ismagilova et al., 2020). Turinio kokybė apima tokius aspektus kaip informacijos tikslumas, aktualumas, gebėjimas įtraukti, suprantamumas bei naujumas (Dabbous & Barakat, 2020; Yeung et al., 2022). Aukštos kokybės turinys, kuris yra informatyvus, aiškiai suprantamas ir aktualus vartotojui, kuria patikimumo įspūdį. Dėl to šiame modelyje turinio vertinimas (per autentiškumo ir kokybės prizmę) veikia kaip tarpininkaujantis veiksnys tarp autorystės ir pasitikėjimo: turinio kilmė – DI ar žmogaus – pirmiausia paveikia tai, kaip vartotojas suvoks turinio autentiškumą ir kokybę, o šie suvokimai nulems, ar vartotojas galiausiai pasitikės turiniu.

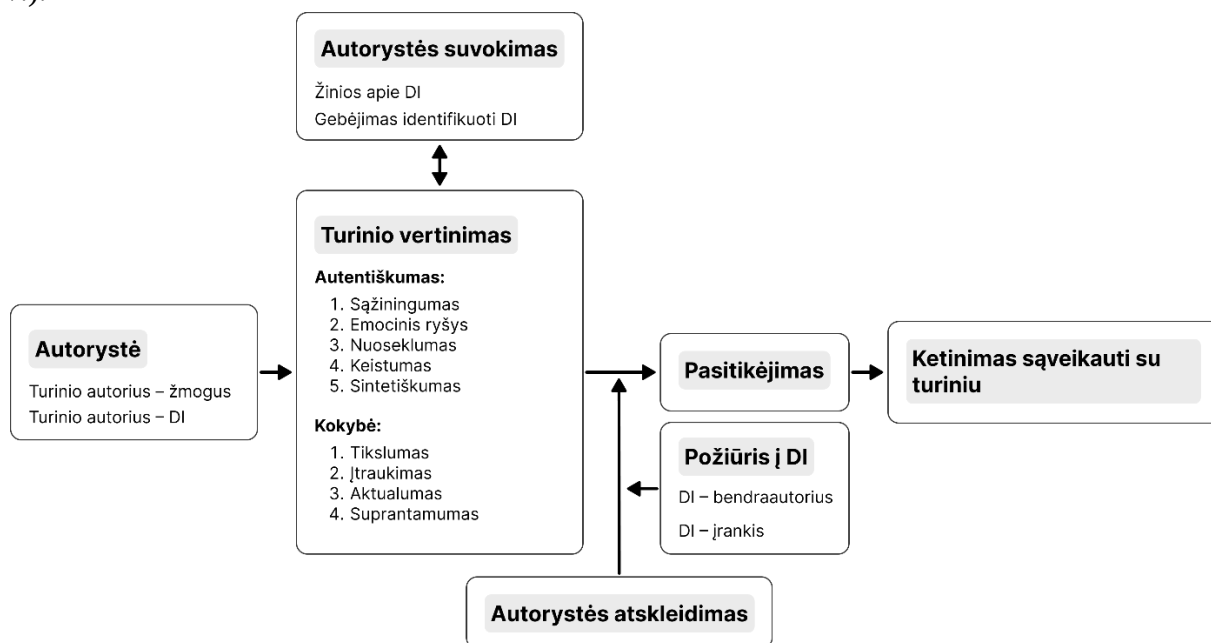
Pasitikėjimas turiniu. Vartotojų pasitikėjimas yra centrinis modelio elementas, jungiantis turinio vertinimą su ketinimais. Atlikti tyrimai patvirtino pasitikėjimo svarbą vartotojų elgsenoje: pasitikėdami pateikiamu turiniu, vartotojai ne tik palankiau jį vertina, bet ir yra labiau linkę įsitraukti – priimti informaciją, ja dalintis su aplinkiniais (Afroogh et al., 2024; Buchanan & Hickman, 2024; Y. Chen et al., 2022; Kim et al., 2015). Pavyzdžiui, jei turinys, kuriuo pasitiki vartotojas, reklamuoja produktą, vartotojas yra linkęs jį įsigyti; jei turinys atrodo patikimas naujienų pranešimas, vartotojas gali juo pasidalinti su kitais. Pasitikėjimo buvimas tampa katalizatoriumi, skatinančiu vartotojų sąveiką su turiniu. Priešingai, pasitikėjimo stoka veikia kaip barjeras – nepasitikėdami turiniu, vartotojai ignoruos arba atmes bet kokią jo daromą įtaką ir vengs su juo sąveikauti (Y. Chen et al., 2022). Svarbu atkreipti dėmesį, jog DI kuriamo turinio atveju pasitikėjimo formavimas gali būti sudėtingesnis: vartotojai neretai jaučia vadinamąjį „technologinį svetimumą“ DI turiniui arba nepasitikėjimą dėl nevisiškai skaidrių DI veikimo procesų. Kitaip tariant, vien žinojimas, kad turinį sukūrė DI, kai kuriems vartotojams sukelia skepticizmą. Vis dėlto, kiti tyrimai pastebi, kad turinio patikimumą vartotojų akyse labiau lemia paties šaltinio vertinimas, pavyzdžiui, kas skelbia turinį – asmuo ar prekių ženklas, nei tai, ar turinys sugeneruotas DI, ar žmogaus. Tai reiškia, kad jei, pavyzdžiui, gerai žinoma ir patikima organizacija dalijasi DI sugeneruotu turiniu, vartotojų pasitikėjimas gali išlikti aukštas dėl pasitikėjimo pačia organizacija.

Požiūris į DI. Ne visi vartotojai vienodai reaguoja į DI kuriamą turinį – čia itin svarbus jų bendras nusiteikimas DI technologijų atžvilgiu. Požiūris į DI apima vartotojo išankstinį pasitikėjimą pačiomis DI technologijomis, jo turimas žinias ir nuostatas apie DI galimybes bei rizikas. Tyrimo kontekste požiūris į DI skirstomas pagal tai, ar vartotojai jį suvokia kaip bendraautorių, ar tik kaip įrankį. Šis požiūris veikia kaip madaruojantis veiksnys, darantis tiesioginę įtaką turinio patikimumo vertinimui, atskleidus turinio autorių. Vartotojai, žvelgiantys į DI kaip į bendraautorių, linkę būti mažiau kritiški DI pagalbai sukurtam turiniui (Sargin, 2024). Priešingai, skeptiškai DI atžvilgiu nusiteikę vartotojai yra linkę nepasitikėti DI generuotu turiniu vien dėl jo kilmės fakto (Kim et al., 2015). Altay'ius ir Gilardi's (2024) pastebi, kad apskritai žmonės iš prigimties dažnai būna skeptiški DI atžvilgiu ir todėl mažiau linkę pasitikėti DI kuriamu turiniu nei žmogaus kurtu. Dėl tokio skirtingo nusistatymo, požiūris į DI modelyje numatytas kaip veiksnys, galintis keisti kitų ryšių stiprumą: palankus požiūris

į DI gali sumenkinti neigiamą DI autorystės atskleidimo įtaką pasitikėjimui turiniu, o nepalankus – priešingai, sustiprinti neigiamą efektą.

Ketinimas sąveikauti su turiniu. Galiausiai, modelio galutinis komponentas yra vartotojo ketinimas sąveikauti su turiniu. Šis ketinimas apima įvairias galimas veikas socialiniuose tinkluose: turinio pasidalijimą, komentavimą, teigiamą ar neigiamą įvertinimą, autoriaus paskyros aplankymą. Ketinimą formuoja ankstesni komponentai, o ypač – pasitikėjimas turiniu. Kai vartotojas pasitiki turiniu, jis yra labiau linkęs atlikti su tuo turiniu susijusius veiksmus (Almira Vania & Daniel Tumpal H., 2023; Li & Yang, 2024; N. N. Mohamed et al., 2023a). Empiriniai tyrimai patvirtina, kad stiprus pasitikėjimas skatina vartotojų išitraukimą: rasta, jog pasitikėjimas informaciją skelbiančiu šaltiniu didina informacijos priėmimą ir galiausiai skatina pirkimo ar naudojimosi ketinimus (Almira Vania & Daniel Tumpal H., 2023; Dabbous & Barakat, 2020; N. N. Mohamed et al., 2023a). Jeigu vartotojas nepasitiki turiniu, tikėtina, kad jis ne tik ignoruos turinio rekomendacijas, bet ir apskritai vengs bet kokios sąveikos – nesidalins, nekomentuos, nespaus „patinka“. Svarbu tai, kad šiame modelyje ketinimai tiesiogiai priklauso nuo pasitikėjimo turiniu: kiti veiksniai – autorystė, autentiškumas, kokybė, požiūris į DI, autorystės atskleidimas – daro netiesioginį poveikį per pasitikėjimą.

Apibendrinant mokslinės literatūros išvalgas, pateikiamas konceptualus modelis, kuris iliustruoja aptartas sąsajas tarp turinio tipo, autentiškumo suvokimo, vartotojų pasitikėjimo ir ketinimų (žr. 3 pav.).



3 pav. Konceptualus dirbtinio intelekto ir žmogaus kuriamo turinio poveikio vartotojų pasitikėjimui ir ketinimams modelis

Modelio pagrindą sudaro keturi pagrindiniai komponentai, sujungti priežastiniais ryšiais: autorystė, turinio vertinimas, pasitikėjimas ir ketinimas sąveikauti. Šių komponentų tarpusavio sąveika formuoja vartotojų elgsenos dinamiką socialiniuose tinkluose. Autorystė modelyje apibrėžiama dviem pagrindiniais tipais: sukurta žmogaus arba dirbtinio intelekto. Turinio vertinimas apima dvi pagrindines kategorijas su specifiniais komponentais: autentiškumas ir kokybiškumas. Šio tyrimo kontekste, turinio autentiškumas vertinamas per jo sąžiningumą, gebėjimą užmegzti emocinį ryšį, nuoseklumą, suvokiamą keistumą ir sintetiškumą. Tuo tarpu kokybė vertinama per turinio tikslumą, gebėjimą įtraukti, aktualumą, suprantamumą. Nuo suvokiamo turinio autoriaus priklauso kaip šis

vertinimas bus atliekamas. Turinio vertinimo elementai tiesiogiai veikia vartotojų pasitikėjimą turiniu. Modelio galutinis rezultatas – vartotojų ketinimas sąveikauti su turiniu, kas apima įvairias elgsenos formas socialiniuose tinkluose: dalijimąsi, komentavimą, reagavimą ar kitokio pobūdžio sąveiką. Paskutinis modelio komponentas atspindi praktinę tyrimo vertę asmenims, kuriantiems įvairaus pobūdžio turinį socialiniuose tinkluose. Konceptualiam modelyje taip pat išskiriami du svarbūs moderatoriai, kurie veikia santykius tarp pagrindinių komponentų: požiūris į DI, t. y. bendras vartotojų nusistatymas DI technologijų atžvilgiu bei autorystės atskleidimas, kurio pagalba siekiama sužinoti kaip kinta vartotojų pasitikėjimas turiniu, jei jo autorius yra žinomas.

Apibendrinant teorinę analizę, išryškėja ženklūs dirbtinio intelekto ir žmogaus kuriamo turinio skirtumai socialiniuose tinkluose, lemiantys vartotojų pasitikėjimą ir elgsenos ketinimus. Dirbtinio intelekto generuojamas turinys pasižymi efektyvumu, sparta ir gebėjimu apdoroti milžiniškus duomenų kiekius, tačiau neretai vertinamas kaip mažiau autentiškas ar asmenišką dėl mechaninio pobūdžio ir ribotų galimybių perteikti asmenines patirtis, kūrybiškumą bei kultūrinius niuansus. Tuo tarpu žmogaus kurtas turinys natūraliai įtraukia emocinį gilumą, empatiją ir individualų kūrybinį braižą, todėl vartotojams atrodo nuoširdesnis bei tikroviškesnis, skatindamas didesnę įsitraukimą ir lojalumą. Dėl šių skirtumų dirbtinio intelekto kuriami pranešimai, nors ir tinkami masinei komunikacijai, dažnai stokoja asmeniškumo ir gali būti vertinami kaip „sintetiški“ ar keliantys „keistumo“ jausmą – tai skatina auditorijos skepticizmą ir mažina pasitikėjimą turiniu. Ne mažiau svarbūs yra turinio autentiškumas ir autorystės atskleidimo skaidrumas. Moksliniai šaltiniai pabrėžia, kad aiškiai nurodytas turinio kūrėjas (žmogus ar DI) mažina auditorijos skepticizmą ir stiprina pasitikėjimą turiniu. Taip yra todėl, kad vartotojai gali lengviau įvertinti šaltinio kompetenciją, patikimumą bei intencijas. Pavyzdžiui, pagal šaltinio patikimumo modelį auditorijos pasitikėjimas informacija priklauso nuo autoriaus patikimumo ir ekspertizmo suvokimo – identifikuotas, gerą reputaciją turintis autorius didina pasitikėjimą, o anoniminis ar neaiškios kilmės turinys kelia abejonių dėl informacijos patikimumo. Vis dėlto autorystės atskleidimas gali sukelti ir atvirkštinį efektą, jei vartotojai skeptiškai žiūri į DI technologijas – turinys, atvirai įvardytas kaip DI generuotas, neretai laikomas mažiau patikimu, ypač kai jam stinga natūralumo ar emocinio ryšio. Apibendrinant, pasitikėjimas išryškėja kaip kritinis veiksnys, lemiantis, ar vartotojai ryšis sąveikauti su turiniu bei įgyvendinti su juo susijusius ketinimus.

Remiantis teorinėmis įžvalgomis, suformuotas konceptualus modelis, kuriame identifikuoti tokie veiksniai kaip autorystės suvokimas, turinio autentiškumas ir kokybė, kaip esminiai vartotojų pasitikėjimą ir ketinimus lemiantys elementai. Šis modelis yra pagrindas tolesnei empirinio tyrimo eigai – kitame skyriuje detalizuojama tyrimo metodologija: apibrėžiamas kokybinio tyrimo metodas, formuojami tyrimo tikslas ir uždaviniai, išskiriamos pagrindinės tyrimo hipotezės, aprašomas tyrimo instrumentas.

3. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams empirinio tyrimo metodologija

Šiame skyriuje pateikiama magistro baigiamojo darbo empirinio tyrimo metodologija: empirinio tyrimo objektas, tikslas, uždaviniai, metodai, tyrimo konstruktas, duomenų analizės procedūros, etika.

3.1. Empirinio tyrimo tikslas, uždaviniai ir hipotezės

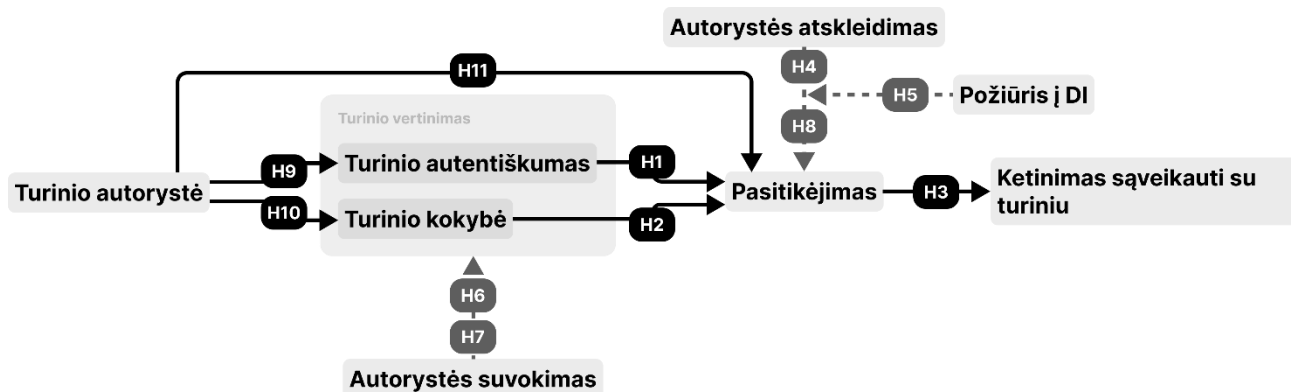
Tyrimo **objektas**: dirbtinio intelekto sprendimų ir žmogaus sukurto turinio socialiniuose tinkluose įtaka vartotojų pasitikėjimui bei ketinimams.

Tyrimo **tikslas**: palyginti dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikį vartotojų pasitikėjimui bei ketinimams.

Tyrimo **uždaviniai**:

1. apibūdinti tyrimo imties sociodemografinės charakteristikas;
2. atskleisti, kokią įtaką vartotojų žinios ir suvokimas apie DI įrankius daro jų gebėjimui tinkamai identifikuoti turinio autorių;
3. išsiaiškinti, kokią įtaką turinio autorystės atskleidimas daro vartotojų pasitikėjimui turiniu;
4. ištirti vartotojų suvokiamo turinio autentiškumo ir kokybės poveikį jų pasitikėjimui turiniu;
5. išanalizuoti, kokią įtaką daro vartotojų pasitikėjimas turiniu jų ketinimams sąveikauti su turiniu;
6. nustatyti kaip vartotojų požiūris į DI moderoja autoriaus atskleidimo poveikį pasitikėjimui turiniu.

Remiantis konceptualių modelių, pateiktu šio darbo 2.5. skyriuje, parengtas tyrimo modelis (žr. 4 pav.). Pastarajame atsispindi ir pagrindinės tyrimo hipotezės.



4 pav. Tyrimo modelis

4 paveiksle pateiktas tyrimo modelis prasideda nuo turinio autorystės – ar jis sukurtas žmogaus ar DI. Šis konstruktas veikia kaip tyrimo pradžios taškas, darantis tiesioginę įtaką turinio vertinimo dimensijoms – autentiškumui ir kokybei, kurie savo ruožtu lemia vartotojų pasitikėjimą turiniu. Autorystės suvokimas vertinamas per vartotojų gebėjimą identifikuoti turinio kilmę, vertinamas pasitelkiant vartotojų žinias ir gebėjimą tinkamai identifikuoti turinio autorių. Autorystės atskleidimas įtraukiamas kaip mediatorius siekiant įvertinti, ar turinio autoriaus identifikavimas reikšmingai veikia vartotojų pasitikėjimą turiniu. Galiausiai, vartotojų ketinimai sąveikauti su turiniu yra sąlygojami jų pasitikėjimo tuo turiniu.

Tyrimo modelyje pažymėtos 11 hipotezių, kurias siekiama patikrinti tyrimo pagalba. H1, H2, H3, H5, H9 ir H10 yra grupinės hipotezės, kurias sudaro keletas sub-hipotezių. **Hipotezių pagrindimas pateiktas 1-ajame priede.** Toliau pateikiamos pirminės tyrimo hipotezės:

H1. Vartotojų suvokiamas turinio autentiškumas teigiamai veikia jų pasitikėjimą turiniu.

H1a. Vartotojų suvokiamas turinio sąžiningumas daro teigiamą įtaką turinio autentiškumo vertinimui.

H1b. Vartotojų suvokiamas turinio emocinis ryšys daro teigiamą įtaką turinio autentiškumo vertinimui.

H1c. Vartotojų suvokiamas turinio nuoseklumas daro teigiamą įtaką turinio autentiškumo vertinimui.

H1d. Vartotojų suvokiamas turinio keistumas daro neigiamą įtaką turinio autentiškumo vertinimui.

H1e. Vartotojų suvokiamas turinio sintetiškumas daro neigiamą įtaką turinio autentiškumo vertinimui.

H2. Vartotojų suvokiama turinio kokybė teigiamai veikia jų pasitikėjimą turiniu.

H2a. Vartotojų suvokiamas turinio tikslumas daro teigiamą įtaką turinio kokybės vertinimui.

H2b. Vartotojų suvokiamas turinio gebėjimas įtraukti daro teigiamą įtaką turinio kokybės vertinimui.

H2c. Vartotojų suvokiamas turinio aktualumas daro teigiamą įtaką turinio kokybės vertinimui.

H2d. Vartotojų suvokiamas turinio suprantamumas daro teigiamą įtaką turinio kokybės vertinimui.

H3. Tikrojo turinio autoriaus atskleidimas moderuoja pasitikėjimą turiniu.

H3a. Dirbtinio intelekto kurto turinio autoriaus atskleidimas daro neigiamą įtaką pasitikėjimui turiniu.

H3b. Žmogaus kurto turinio autoriaus atskleidimas daro teigiamą įtaką pasitikėjimui turiniu.

H4. Vartotojų pasitikėjimo turiniu pokytis teigiamai veikia jų polinkį sąveikauti su turiniu.

H5. Požiūris į DI moderuoja kaip vartotojų pasitikėjimo pokytį veikia tikrojo turinio autoriaus atskleidimas.

H5a. Dirbtinio intelekto turinio tikrojo autoriaus atskleidimas neigiamai veikia vartotojų, kurie į DI žvelgia kaip į bendraautorius, pasitikėjimo turiniu pokyčiui.

H5b. Dirbtinio intelekto turinio tikrojo autoriaus atskleidimas teigiamai veikia vartotojų, kurie į DI žvelgia kaip į įrankį, pasitikėjimo turiniu pokyčiui.

H6. Vartotojų žinios apie DI įrankius teigiamai veikia jų gebėjimą identifikuoti tikrąjį turinio autorių.

- H7.** Vartotojų suvokiamas gebėjimas tinkamai identifikuoti DI turinį teigiamai veikia jų gebėjimą identifikuoti DI turinį.
- H8.** Vartotojų suvokiamas pasitikėjimo turiniu didėjimas žinant autorių teigiamai veikia jų pasitikėjimą turiniu, atskleidus tikrąjį turinio autorių.
- H9.** Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio autentiškumui nei dirbtinio intelekto sukurtas turinys.
- H9a.** Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio sąžiningumui nei dirbtinio intelekto sukurtas turinys.
- H9b.** Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio emociniam ryšiui nei dirbtinio intelekto sukurtas turinys.
- H9c.** Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio nuoseklumui nei dirbtinio intelekto sukurtas turinys.
- H9d.** Žmogaus sukurtas turinys daro stipresnę neigiamą poveikį vartotojų suvokiamam turinio keistumui nei dirbtinio intelekto sukurtas turinys.
- H9e.** Žmogaus sukurtas turinys daro stipresnę neigiamą poveikį vartotojų suvokiamam turinio sintetiškumui nei dirbtinio intelekto sukurtas turinys.
- H10.** Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamai turinio kokybei nei dirbtinio intelekto sukurtas turinys.
- H10a.** Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio tikslumui nei dirbtinio intelekto sukurtas turinys.
- H10b.** Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio gebėjimui įtraukti nei dirbtinio intelekto sukurtas turinys.
- H10c.** Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio aktualumui nei dirbtinio intelekto sukurtas turinys.
- H10d.** Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio suprantamumui nei dirbtinio intelekto sukurtas turinys.
- H11.** Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų pasitikėjimui turiniu nei dirbtinio intelekto sukurtas turinys.

3.2. Empirinio tyrimo tipas, duomenų rinkimo metodas ir tyrimo konstrukto operacionalus apibūdinimas

Tyrimui pasirinkta pozityvistinė epistemologinė pozicija, kadangi siekiama atrasti konkrečius veiksnius, darančius įtaką vartotojų elgsenai bei ketinimams per objektyvią duomenų analizę. Tokia pozicija padės identifikuoti dėsningumus bei priežasties-pasekmės ryšius.

Tyrimui pasirinktas **kiekybinis tyrimas**. Kiekybiniai tyrimai ypač tinka analizuojant sudėtingą žmogaus elgseną, socialinę dinamiką ir kultūrinius kontekstus pasitelkiant kiekybiškai išmatuojamus ir susisteminamus elgsenos aspektus (Jamieson et al., 2023).

Tyrimo tikslui pasiekti naudojamas scenarijais grįstas kiekybinis eksperimentinis metodas. Šis tyrimo metodas yra itin veiksmingas, analizuojant tiriamųjų suvokiamą pasitikėjimą ir ketinimą veikti turinio atžvilgiu. Šio tyrimo kontekste tai būtų skirtingi turinio socialiniuose tinkluose autoriai: dirbtinis intelektas ir žmogus.

Atsižvelgiant į tyrimo tikslą, turimus išteklius ir tyrimo laikotarpį, tyrimo duomenims surinkti pasirinkta **internetinė anketinė apklausa**.

Atkreipiant dėmesį į tai, jog vartotojų kultūrinis ir socialinis kontekstas daro itin reikšmingą įtaką turinio suvokimui, pasirinkta analizuoti konkrečios šalies tendencijas (Guan et al., 2022; Krishen et al., 2021). Tyrimo **populiacija** apibrėžiama kaip pilnamečiai Lietuvos gyventojai, turintys ne mažesnę kaip vidurinį išsilavinimą, kalbantys lietuvių kalba ir aktyviai naudojantys socialinius tinklus.

Siekiant surinkti reprezentatyvų kiekį respondentų **pasirinkta patogioji atranka** duomenims rinkti. Patogioji atranka leidžia greičiau pasiekti respondentus ir surinkti duomenis per trumpesnę laiką, ypač kai tyrimo resursai ir laikas yra riboti (Jamieson et al., 2023).

Remiantis tyrimo modeliu (žr. 4 pav.) ir išsikeltomis hipotezėmis, buvo sudaryta apklausa lietuvių kalba. Šis sprendimas buvo priimtas siekiant patikrinti hipotezes Lietuvoje gyvenančių lietuvių mastu.

Tyrimo anketa (žr. 2 priedą) buvo sudaryta pagal 5 paveiksle pateiktą schemą.



5 pav. Anketos schema

Remiantis jau esamais moksliniais tyrimais (Almira Vania & Daniel Tumpal H., 2023; Belanche et al., 2024; Gu et al., 2024; Sargin, 2024), tyrimo dalyviams anketos pradžioje pateikiama tyrimo santrauka ir pagrindiniai tikslai. Tyrimo modelyje (žr. 4 pav.) išskirti konstruktai anketoje nėra atskleidžiami, taip siekiant užtikrinti tyrimo objektyvumą ir išpildyti siekti įvertinti vartotojų pasitikėjimą ir ketinimus turinio atžvilgiu.

Respondentai, paspaudę bendrą tyrimo nuorodą, automatiškai buvo nukreipiami į vieną iš dviejų anketų rotacijos principu – tam panaudota Linkly „rotation“ funkcija. Tokiu būdu užtikrinta, kad abiem scenarijams surenkama pakankamai atsakymų ir grupės būtų panašaus dydžio.

Po įžangos, tyrimo dalyviai kviečiami atsakyti į keletą klausimų, susijusių su socialinių tinklų naudojimo įpročiais, kad būtų galima įvertinti respondento patirtį su socialiniais tinklais (Almira Vania & Daniel Tumpal H., 2023; Gu et al., 2024; Ostic et al., 2021; Sargin, 2024). Socialinių tinklų naudojimo įpročių klausimai taip pat atlieka kontrolės vaidmenį – asmenys, kurie nesinaudoja socialiniais tinklais, arba tai daro labai retai, yra nukreipiami į apklausos pabaigą ir jų atsakymai tolimesniame tyrime nėra nagrinėjami.

Toliau respondentai kviečiami įvertinti savo gebėjimą atpažinti dirbtinio intelekto turinį ir įvertinti turinio autoriaus svarbą nustatant turinio patikimumą (Y. Chen et al., 2022; Sargin, 2024). Šie klausimai užduodami būtent šiame tyrimo etape siekiant paskatinti žmonės patiems įsivertinti gebėjimą atskirti DI turinį nuo žmogaus. Tolimesnė šio tyrimo dalis šį įsivertinimą patikrina praktiškai prašydama identifikuoti turinio autorių (Park et al., 2024).

Toliau seka klausimai, kurių pagrindu nustatomas respondentų požiūris į dirbtinį intelektą (Sargin, 2024). Ši informacija itin naudinga siekiant sužinoti kokią įtaką turinio vertinimui daro išankstinis nusistatymas už ar prieš dirbtinį intelektą ir jo siūlomus sprendimus.

Tuomet respondentai nukreipiami į vieną iš dviejų scenarijų:

1. Į pirmąją respondentų grupę nukreipti tyrimo dalyviai supažindinami su dirbtinio intelekto pagalba sugeneruotu tipiniu socialinių tinklų įrašu humanoidinių robotų tema. Įrašo tekstą sugeneravo „ChatGPT 4o“, o nuotrauką – „Google Gemini Advanced 2.0 Flash“. Turiniui sugeneruoti naudojamas užduotis pateikiamos 3-ajame darbo priede.
2. Į antrąją grupę patekusiems respondentams parodomas „Greater Zurich Area“ projektų ir bendruomenės vadovo Noah'o Zahnd'o „LinkedIn“ įrašas humanoidinių robotų tema, išleistas 2025-ųjų kovo 22 dieną.

Atkreipiant dėmesį į tai, jog siekiama ištirti Lietuvoje gyvenančių lietuviškai kalbančių piliečių nuomonę, vartotojams buvo suteikta galimybė pamatyti tiek anglišką, tiek lietuvišką teksto variantą, imituojant socialinių tinklų vertimo paslaugą ir užtikrinant, kad kalbos barjeras nedarytų įtakos respondento galimybei dalyvauti tyrime. Socialinių tinklų įrašų tekstas buvo išverstas mašininio būdu, siekiant imituoti realių socialinių tinklų funkcionalumą. Zahnd'o įrašo vertimas adaptuotas tiesiai iš „LinkedIn“, tuo tarpu DI pagalba sukurtas įrašas buvo išverstas su „ChatGPT 4o“ pagalba. Socialinių tinklų įrašą išversti į lietuvių kalbą pasirinko 83 proc. respondentų.

Vengiant šališkumo, konkrečios įmonės ar asmenys anketoje nebuvo minimi. Respondentai, atpažinę tyrime minimus juridinius ir fizinius asmenis gali būti linkę į klausimus reaguoti subjektyvumo pagrindu. Respondentai iš pradžių nematė, kas yra turinio autorius – ši informacija atskleista vėliau apklausoje, kad būtų galima įvertinti autorystės atskleidimo efektą.

Siekiant sudominti respondentus, pasirinkta šiandieną itin aktuali ir emocinį įsitraukimą skatinanti tema apie humanoidinius robotus. Humanoidinių robotų kūrimas ir diegimas – viena greičiausiai besivystančių sričių dirbtinio intelekto ir robotikos pasaulyje. Jie tampa ne tik mokslinės fantastikos objektu, bet realia dalimi verslo, paslaugų ir net kasdienio gyvenimo. Pasaulyje pastebimas augantis humanoidų diegimas sveikatos priežiūros, švietimo, aptarnavimo sektoriuose – o tai rodo, kad jie tampa plačiai pritaikomi. Humanoidiniai robotai kelia daug klausimų apie darbo rinkos pokyčius, žmogaus ir mašinos santykį, socialinę bei emocinę integraciją. Diskusijos apie robotų etiką, teises, bei jų vietą visuomenėje tampa vis garsesnės ir labiau įtraukiančios plačiąją auditoriją.

Pirmajame scenarijuje pateikto dirbtinio intelekto pagalba sugeneruoto įrašo ilgis siekė 815 spaudos ženklų (be tarpų) anglų kalba bei 803 spaudos ženklai lietuvių kalba; tuo tarpu žmogaus sukurto turinio tekstą sudarė 862 spaudos ženklai (be tarpų) anglų kalba ir 1000 spaudos ženklų lietuvių kalba.

Po įrašo peržiūros sekė klausimų blokas pirminiam turinio vertinimui – čia dalyviai įvertino turinio autentiškumą, turinio kokybę ir pirminį pasitikėjimą turiniu pagal pateiktas skales (Almira Vania & Daniel Tumpal H., 2023; Belanche et al., 2024; Gu et al., 2024; Huschens et al., 2023; Sargin, 2024). Pateiktus teiginius vartotojai turėjo įvertinti nuo 1 – visiškai nesutinku, iki 5 – visiškai sutinku.

Autoriaus atskleidimas: kitoje sekcijoje respondentui atskleidžiama, kas iš tikrųjų yra įrašo autorius – DI arba žmogus – ir paklausama, ar pasikeitė jo pasitikėjimas turiniu sužinojus tikrąjį turinio autorių. Taip tiesiogiai matuojamas autorystės atskleidimo poveikis pasitikėjimui.

Toliau respondentams pateikiami klausimai apie ketinimus sąveikauti su turiniu, pavyzdžiui, ar respondentas būtų linkęs pamėgti, komentuoti, pasidalinti tokiu įrašu, ar domėtusi plačiau (Dabbous & Barakat, 2020; Gu et al., 2024; Park et al., 2024). Šia klausimyno dalimi siekiama išsiaiškinti kokią įtaką pasitikėjimas turiniu daro vartotojų ketinimui su juo sąveikauti.

Galiausiai, remiantis Ostic et al. (2021), Gu et al. (2024) bei Huschens et al. (2023) atliktais tyrimais, paskutinėje anketos dalyje tyrimo dalyviai yra kviečiami pasidalinti savo sociodemografinę informaciją: lytimi, amžiumi, išsilavinimu patirtimi su LLM (angl. *large language model*) (žr. 2 priedą).

Visi tyrimo konstruktai ir matavimo skalės pateikiami 13-oje lentelėje.

13 lentelė. Tyrimo konstruktai ir matavimo skalės

Kategorija		Teiginys	Literatūros šaltiniai
Autorystė		1. Aš galiu lengvai atpažinti, ar socialinio tinklo turinys buvo sukurtas DI ar žmogaus. 2. Jei žinau, kas yra turinio autorius, esu linkęs labiau juo pasitikėti.	Adaptuota pagal Chen et al. (2024), Sargin (2024) bei Altay & Gilardi (2024b)
Požiūris į DI	DI kaip bendraautorius	1. Žvelgiu į DI kaip į bendraautorį, padedantį perteikti informaciją 2. Žmogus, pasitelkdamas DI įrankius, kuria vertingiausius įrašus socialiniuose tinkluose 3. Manau, jog DI yra lygiavertis partneris, kuriant turinį	Adaptuota pagal Sargin (2024)
	DI kaip įrankis	3. DI yra tik įrankis, kuriuo naudojasi žmogus 4. Manau, kad žmogus turi tiesiogiai kontroliuoti, kokį turinį kuria DI	Adaptuota pagal Sargin (2024)

Autentiškumas	Sąžiningumas	4. Įrašas atrodo nuoširdus 5. Įrašė manipuliuojama statistiniais faktais (r) 6. Įrašė nėra faktinių klaidų	Adaptuota Belanche et al. (2024), Almira Vania & Daniel Tumpal H. (2023)
	Emocinis ryšys	7. Įrašė išvelgiu tikrosios žmogiškosios patirties 8. Įrašas demonstruoja gyvybingumą ir turi savą asmenybę	Adaptuota pagal Buzeta et al. (2020), Gu et al. (2024)
	Nuoseklumas	9. Įrašas pasižymi nuoseklumu 10. Įrašo formatas ir struktūra yra nuspėjami	Adaptuota pagal Sargin (2024) ir Gu et al. (2024)
	Keistumas	11. Įrašas yra šiurpus (r) 12. Įrašas yra neįprastas (r)	Adaptuota pagal Gu et al. (2024)
	Sintetiškumas	13. Įrašas, kaip vaizdo ir teksto visuma, sudaro įspūdį, jog yra sudėliotas iš nesusijusių dedamųjų (r) 14. Kai kurios įrašo detalės atrodo nenatūraliai (r) 15. Įrašas pasižymi sintetiškumu (r) 16. Įrašas, kaip visuma, sudaro nesuderintų derinių įspūdį (r)	Adaptuota pagal Gu et al. (2024)
Kokybė	Tikslumas	17. Įrašas atrodo faktiškai teisingas 18. Papildomai tikrinčiau pagrindinius teiginius, pateiktus įrašė (r)	Adaptuota pagal Huschens et al. (2023)
	Įtraukimas	19. Man buvo lengva išlaikyti dėmesį skaitant įrašą 20. Įrašą peržvelgiau nuo pradžios iki pabaigos 21. Įrašas yra įdomus	Adaptuota pagal Huschens et al. (2023)
	Aktualumas	22. Įrašas atkreipia dėmesį į kultūrinį kontekstą 23. Įrašas orientuotas į temas, kurios aktualios dabar 24. Įrašas turi išliekamosios vertės 25. Įrašė pastebėjau dar negirdėtų minčių	Adaptuota pagal Belanche et al. (2024), Gu et al. (2024)
	Suprantamumas	26. Įrašą buvo lengva suprasti 27. Įrašė pateikiami palyginimai palengvina jo supratimą 28. Įrašas pateikiamas logiškai ir struktūruotai	Adaptuota pagal Huschens et al. (2023)
Pasitikėjimas		29. Įrašo autoriumi galima pasitikėti 30. Įrašas yra patikimas 31. Įrašo autorius yra žmogus	Adaptuota pagal Huschens et al. (2023), Almira Vania & Daniel Tumpal H. (2023) ir Park et al. (2024)
Ketinimai turinio atžvilgiu		32. Tikėtina, jog perskaityčiau šį įrašą socialiniuose tinkluose 33. Tikėtina, jog apsilankyčiau įrašo autoriaus paskyroje 34. Tikėtina, jog pasidalinčiau šiuo įrašu su draugais socialiniuose tinkluose 35. Tikėtina, jog sureaguočiau į šį turinį paspausdama (-s) „patinka“	Adaptuota pagal Dabbous & Barakat (2020), Park et al. (2024), Gu et al. (2024)

3.3. Empirinio tyrimo imties atrankos procedūros, tyrimo etika ir analizės metodai

Tinkamos tyrimo imties nustatymas yra itin svarbus, siekiant užtikrinti statistinį pagrįstumą, kartu atsižvelgiant į praktinius tyrimo apribojimus. Siekiant tinkamai nustatyti tyrimo imtį, buvo atsižvelgta į keletą dedamųjų (Das et al., 2016).

- 1. Patikimumo lygmuo: dauguma akademinių tyrimų dirbtinio intelekto ir žmogaus kuriamo turinio tema pasirenka 95 proc. patikimumo lygmenį (Belanche et al., 2024; Das et al., 2016; Park et al., 2024; Sargin, 2024; M. Yang & Rau, 2024).
- 2. Paklaidos riba: kuo mažesnė paklaidos riba, tuo didesnės apimties reikia. Remiantis mokslininkų tyrimais, dažniausiai pasirenkama paklaidos riba yra 5 proc. (Belanche et al., 2024; Das et al., 2016; Park et al., 2024; Sargin, 2024; M. Yang & Rau, 2024).
- 3. Populiacijos dydis: remiantis Oficialiosios statistikos portalo duomenimis (2024), šiuo metu Lietuvoje gyvena 2 889 402 asmenys iš kurių 82,5 proc. yra lietuviai. Iš Lietuvoje gyvenančių kitataučių, Oficialiosios statistikos portalo duomenimis (2021), vidutiniškai 9,9 proc. moka lietuvių kalbą. Tuo tarpu asmenys jaunesni nei 18 metų sudaro vidutiniškai 15 proc. Lietuvos gyventojų. Remiantis Bitės profų tyrimu (2024), 87 proc. suaugusiųjų Lietuvoje naudoja socialinius tinklus. Šiuo atveju tyrimo populiaciją sudaro 1 762 988 asmuo.
- 4. Kadangi tyrimui atlikti naudojami du skirtingi scenarijai, atsakymų pasiskirstymas turėtų sudaryti santykį 1:1, t.y. 50 proc. respondentų turėtų atsakyti į klausimus pirmajame scenarijuje, o kiti 50 proc. – į kitą.

Pasinaudojus imties dydžio skaičiuokle „PureCalculations“ ir įvedus aukščiau įvardintus duomenis, nustatyta, jog reprezentatyviam tyrimui atlikti reikia 385 respondentų atsakymų. Atkreipiant į tai, jog testuojami skirtingi scenarijai, į kiekvieną jų turėtų įvertinti 193 asmenys. 14-oje lentelėje išskiriami kiti tyrimai panašia tema ir jų imties dydžiai.

14 lentelė. Imties dydis moksliniuose darbuose

Autorius	Tyrimo pavadinimas	Imties dydis
Lappeman et al. (2023)	Trust and digital privacy: willingness to disclose personal information to banking chatbot services	726
Campos et al. (2023)	Revisiting the debate: documenting biodiversity in the age of digital and artificially generated images	511
Graefe & Bohlken (2020)	Automated journalism: A meta-analysis of readers' perceptions of human-written in comparison to automated news	388
Scharowski et al. (2023)	Exploring the effects of human-centered AI explanations on trust and reliance	380
Belanche et al. (2024)	Human versus virtual influences, a comparative study	340
M. Yang & Rau (2024)	Trust building with artificial intelligence: comparing with human in investment behaviour, emotional arousal and neuro activities	320
Park et al. (2024)	AI vs. human-generated content and accounts on Instagram: User preferences, evaluations, and ethical considerations	251
Vidurkis		410

Mokslininkai, analizuodami dirbtinio intelekto ir žmogaus kuriamo turinio įtaką vartotojų pasitikėjimui ir ketinimams, vidutiniškai apklausė 410 respondentų. Šis skaičius yra labai artimas apskaičiuotam imties dydžiui. Įvertinus tiek apskaičiuotą imties dydį, tiek remiantis kitais tyrimais, nustatyta tyrimo imtis: 398 respondentai, 199 kiekvienam scenarijui.

Prieš paskelbiant anketą, apklausa buvo pasidalinta su tyrimo populiacijos aprašymą atitinkančiais asmenimis, taip atliekant pilotinį tyrimą. Į šį bandomąjį tyrimą įsitraukė 11 asmenų. Pilotinio tyrimo tikslas buvo įvertinti anketos aktualumą, tinkamumą ir potencialius trūkumus. Taip pat šis tyrimas leido nustatyti apklausos užpildymo trukmę – vidutiniškai respondentai apklausą užpildė per 9 min. 58 sek. Vis dėlto, įvertinus tai, jog realaus tyrimo metu anketą pildys daug daugiau asmenų, tyrimo aprašyme nuspręsta laiko režius pildymui įvardyti kaip nuo 10 iki 12 min.

Faktinė tyrimo imtis: anketą užpildė 466 respondentai. Tolimesnei statistinei analizei buvo atrinktos 461 anketos, kadangi 3 respondentai socialiniais tinklais naudojami rečiau nei kartą per mėnesį, o 2 – jais visai nesinaudoja. Pašalinus šiuos atvejus, tolesnei analizei buvo naudojamos 232 anketos, kuriose respondentai matė pirmojo scenarijaus (socialinių tinklų įrašas sukurtas dirbtinio intelekto pagalba) variantą, ir 229 anketos, susijusios su antruoju scenarijumi (įrašas sukurtas žmogaus).

Tyrimo etika. Anketos dalyviams, prieš įsitraukiant į tyrimą, pateikiama išsami informacija apie jo tikslą, eigą, apklausos trukmę bei klausimų pobūdį. Pabrėžiama, jog gaunami duomenys bus naudojami apibendrintai ir tik magistro baigiamojo darbo uždaviniams įgyvendinti. Užtikrinant respondentų anonimiškumą, anketoje neprašoma informacijos, pagal kurią būtų įmanoma identifikuoti asmenį. Tyrimo dalyviai turi galimybę bet kuriuo metu nutraukti dalyvavimą tyrime, o iki tol jų pateikta informacija būtų sunaikinta. Tyrimo respondentams suteikiama galimybė savarankiškai pareikšti norą susipažinti su tyrimo rezultatais. Dalyvavimas tyrime grindžiamas savanoriškumo principu. Tyrimui vykdyti, duomenims surinkti ir analizuoti pasitelkiamos saugūs ir akademinėje bendruomenėje gerą reputaciją turintys įrankiai.

Anketinė apklausa buvo parengta naudojantis „Tally.com“ platforma, o jos platinimui pasitelktas „Link.ly“ įrankis, turintis „sukimo“ (angl. *rotation*) funkciją. Ši funkcija leidžia paeiliui nukreipti respondentus į skirtingas anketos versijas, taip užtikrinant, kad abu scenarijai sulauktų vienodo respondentų skaičiaus. Tyrimo anketa buvo platinama socialiniuose tinkluose „Facebook“ ir „LinkedIn“ iš asmeninės paskyros, patalpinus įrašą apie vykdomą tyrimą (žr. 4 priedą). Taip pat įrašas apie vykdomą tyrimą buvo paskelbtas ir specializuotose socialinių tinklų grupėse, tokiose kaip „Apklausa!“ ir „APKLAUSOS“. Reikiamą respondentų skaičių pavyko pasiekti per 7 dienų laikotarpį: 2025 m. balandžio 2-8 dienomis (imtinai).

Analizės metodai. Surinkti duomenys analizuojami naudojant „IBM SPSS Statistics“ programinę įrangą; atliekama aprašomoji statistinė analizė: imties sociodemografinio pasiskirstymo rodikliai, kiekvieno tiriamojo kintamojo rezultatų pasiskirstymai – vidurkiai, standartiniai nuokrypiai, minimumai ir maksimumai; įvertintas duomenų kokybiškumas: atliekama eksploracinė faktorinė analizė, siekiant patikrinti skalių patikimumą; patikrinama, ar klausimyno teiginiai grupuojasi į numatytus veiksnius (konstruktus); įvertinamas kiekvienos skalės vidinis suderinamumas: apskaičiuoti Cronbach'o alfa koeficientai; apskaičiuojamos koreliacijos tarp pagrindinių kiekybinių kintamųjų; įvertinami skirtumai tarp dviejų eksperimentinių grupių ir atitinkamų veiksmų sąveika; atliekamas Mann-Whitney U testas, atliekama paprastoji ir daugialypė regresinė analizė; atliekama moderacijos analizė; duomenų normalaus skirstinio vertinimas; chi-kvadrato testas.

4. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams empirinį tyrimo rezultatai

Šiame skyriuje pristatomi ir interpretuojami empirinio tyrimo rezultatai. Aptariamos tyrimo išvados, pasiūlymai tolimesniems tyrimams, išryškinami tyrimo apribojimai bei formuluojami pasiūlymai.

4.1. Bendri respondentų demografiniai ir sociografiniai rodikliai

Sociodemografiniai tyrimo dalyvių duomenys, tokie kaip lytis, amžius, išsilavinimas bei patirtis tiek su dirbtiniu intelektu grįstais įrankiais, tiek su socialiniais tinklais, gali turėti reikšmingos įtakos tyrimo rezultatų interpretacijai konteksto požiūriu. Dėl šios priežasties pirmiausia buvo atlikta nuodugni respondentų demografinių charakteristikų analizė. Apibendrinti rezultatai pateikiami 15 lentelėje.

15 lentelė. Respondentų lyties ir išsilavinimo pasiskirstymas (N=461)

Charakteristika	Reikšmė	N	Procentinė dalis
Lytis	Vyras	204	44,3 %
	Moteris	256	55,5 %
	Kita	1	0,2 %
Išsilavinimas	Pagrindinis	30	6,5 %
	Vidurinis	47	10,2 %
	Profesinis	40	8,7 %
	Aukštasis (neuniversitetinis)	119	25,8 %
	Aukštasis (universitetinis)	225	48,8 %
DI įrankių naudojimo dažnumas	Kasdien	70	15,2 %
	Keletą kartų per savaitę	118	25,6 %
	Kartą per savaitę	54	11,7 %
	Kartą per mėnesį	65	14,1 %
	Dar rečiau	74	15,8 %
	Niekada	43	9,3 %
	Nesu susipažinęs (-usi) su šiomis programomis	38	8,2 %

Iš 15-oje lentelėje pateiktų duomenų matyti, kad tyrime dalyvavo daugiau moterų nei vyrų – atitinkamai 55,5 proc. ir 44,3 proc. Nedidelė dalis respondentų (0,2 proc.) nurodė kitą lyties tapatybę. Vertinant išsilavinimo pasiskirstymą, didžioji dalis apklaustųjų buvo įgiję aukštąjį universitetinį (48,8 proc.) arba neuniversitetinį (25,8 proc.) išsilavinimą. Profesinį pasirinko 8,7 proc., vidurinį – 10,2 proc., o pagrindinį išsilavinimą turėjo 6,5 proc. tyrimo dalyvių.

Apžvelgiant dirbtinio intelekto įrankių naudojimo dažnumą, daugiausia respondentų – 25,6 proc. – nurodė, kad jais naudojasi keletą kartų per savaitę. Kasdien DI įrankius naudoja 15,2 proc., o kartą per savaitę – 11,7 proc. Tyrimo duomenys rodo, kad 14,1 proc. respondentų naudoja DI kartą per mėnesį, o dar rečiau – 15,8 proc. Tuo tarpu visiškai nenaudojančių DI įrankių dalis siekia 9,3 proc. Taip pat 8,2 proc. respondentų pažymėjo, kad nėra susipažinę su tokiomis programomis.

Respondentų amžiaus santykinė skalė buvo perkoduota į keturias amžiaus grupės. Siekiant išskirti amžiaus grupes, kurios būtų apimtų kiek įmanoma vienodesnį respondentų skaičių, amžiaus kintamojo duomenys buvo suskirstyti kvartiliais. Remiantis kvartilių pagalba išskirtais amžiaus rėžiais, buvo sukurtas naujas kintamasis su amžiaus grupėmis naudojantis *Transform > Compute Variable* funkcija. Bendros respondentų amžiaus tendencijos pateikiamos 16 lentelėje.

16 lentelė. Respondentų pasiskirstymas pagal amžių (N=461)

Minimali reikšmė	Maksimali reikšmė	Vidurkis	Mediana	Standartinis nuokrypis
18	74	36,25	35	11,811
Amžiaus grupė		N	Procentinė dalis	
<= 26 m.		117	25,4 %	
27 – 35 m.		117	25,4 %	
36 – 44 m.		123	26,7 %	
>= 45 m.		104	22,6 %	

Tyrime dalyvavusių asmenų amžius svyravo nuo 18 iki 74 metų. Apskaičiuota, kad dalyvių vidutinis amžius buvo 36,25 metų, o mediana – 35 metai. Standartinis nuokrypis siekė 11,811, nurodydamas gan didelę amžiaus įvairovę tarp respondentų. Daugiausiai respondentų priklausė 36–44 metų amžiaus grupei – šiai kategorijai atiteko 26,7 proc. visų tyrime dalyvavusių asmenų. Po lygiai – po 25,4 proc. – sudarė tiek iki 26 metų, tiek 27–35 metų amžiaus grupės. Mažiausiai dalyvių buvo 45 metų ir vyresnių grupėje – ši grupė sudarė 22,6 proc. visų respondentų.

Respondentų socialinių tinklų naudojimo tendencijos tyrimo kontekste apibrėžiamos kaip socialinių tinklų naudojimo dažnumas, trukmė bei socialinių tinklų platformos, kurias naudoja respondentai. Žemiau pateikta susisteminta informacija, apibūdinanti šiuos aspektus (žr. 17 lentelę).

17 lentelė. Respondentų socialinių tinklų naudojimo tendencijos (N=461)

Charakteristika	Reikšmė	N	Procentinė dalis
Socialinių tinklų naudojimo dažnumas	Keletą kartų per dieną	292	63,3 %
	Kartą per dieną	58	12,6 %
	Keletą kartų per savaitę	46	10 %
	Kartą per savaitę	22	4,8 %
	Keletą kartų per mėnesį	43	9,3 %
Socialinių tinklų naudojimo trukmė	Mažiau nei 4 val.	224	48,6 %
	4-8 val.	141	30,6 %
	9-12 val.	67	14,5 %
	Daugiau nei 12 val.	29	6,3 %
Socialiniai tinklai	Facebook	300	65,1 %
	Instagram	275	59,7 %
	WhatsApp	149	32,3 %
	X	101	21,9 %
	Pinterest	170	36,9 %
	YouTube	257	55,7 %

	LinkedIn	151	32,8 %
	TikTok	159	34,5 %
	Kita	7	1,5 %

Dauguma respondentų (63,3 proc.) socialinius tinklus naudoja kelis kartus per dieną, 12,6 proc. juos tikrina bent kartą per dieną. Rečiau, t. y. keletą kartų per savaitę ar dar retesniu dažnumu, socialinius tinklus naudoja apie ketvirtadalis apklaustųjų. Kalbant apie laiką, skiriamą socialiniams tinklams, beveik pusė (48,6 proc.) respondentų nurodė praleidžiantys mažiau nei 4 valandas per dieną prie socialinių tinklų, tuo tarpu 6,3 proc. pažymėjo, kad socialiniuose tinkluose praleidžia daugiau nei 12 valandų. Populiariausi socialiniai tinklai tarp tyrimo dalyvių buvo „Facebook“ ir „Instagram“, atitinkamai naudojami 65,1 proc. ir 59,7 proc. respondentų. Mažiausiai respondentų naudoja tokias platformas kaip WhatsApp (32,3 proc.) ir X (21,9 proc.). Visi asmenys, kurie pasirinko parinktį „Kita“, papildomai nurodė naudojantys „Reddit“ platformą.

4.2. Tyrimo instrumento patikimumo ir kokybės įvertinimas

Tiriamųjų konstrukto struktūra buvo analizuojama taikant tiriamąją faktorinę analizę (žr. 6 priedą), siekiant įvertinti skalių tinkamumą. Faktorinė analizė buvo vykdoma dviem etapais: pirmiausia vertintos skalės, laikomos veiksniais, o vėliau – skalės, laikomos pasekmėmis. Abiem etapais buvo naudojamas pagrindinių komponentų išskyrimo metodas. Tokia analizės struktūra pasirinkta todėl, kad dauguma skalių buvo sudarytos adaptuojant jas iš įvairių mokslinių šaltinių. Sprendimui taip pat turėjo įtakos tai, jog tiriamą tema šiuo metu yra palyginti nauja, o mokslinėje literatūroje stokojama šaltinių, kurie vienodai atsižvelgtų į tiek tekstinę, tiek vaizdinę informaciją socialinių tinklų įrašų kontekste. Toks metodinis pasirinkimas taip pat leidžia įvertinti skalių, sudarytų tik iš dviejų teiginių, tinkamumą. Tokiais atvejais įprastinė faktorinė analizė būtų negalima dėl per žemo, t. y. mažesnio nei 0,500, KMO (Kaiser-Meyer-Olkin) imties adekvatumo indekso (Almquist et al., 2019). Prieš atliekant faktorinę analizę, visi reversiniai teiginiai (žr. 14 lentelę) buvo perkoduoti, naudojant *Transform > Compute Variable* funkciją. Duomenų tinkamumas faktorių analizei buvo vertinamas taikant KMO indeksą ir Bartleto sferiškumo testą, kurių reikšmės atsispindi 18-oje lentelėje.

18 lentelė. Skalių tinkamumo vertinimas (N=461)

Etapas	Kintamasis	Dimensija	Teiginių skaičius	KMO	Bartlet'o sferiškumo rodiklis
Veiksniai	Požiūris į DI	DI kaip bendraautorius	3	0,851	< 0,001
		DI kaip įrankis	2		
	Autentiškumas	Sąžiningumas	3		
		Emocinis ryšys	2		
		Nuoseklumas	2		
		Keistumas	2		
		Sintetiškumas	4		
	Kokybė	Tikslumas	2		
		Įtraukimas	3		
		Aktualumas	4		
		Suprantamumas	2		

	Pasitikėjimas	2		
Pasekmės	Elgsena	4	0,695	< 0,001

Pasak Piligrimienės (2016), jei KMO mato reikšmė maža, tai nagrinėjamų kintamųjų faktorinę analizę nerezultatyvi. Vadovaujantis autorės pateikta KMO reikšmių interpretacijų lentelė, faktorinę analizę galima atlikti, jei KMO rodiklis yra didesnis nei 0,5. 19-oje lentelėje pateikti duomenys rodo, jog veiksmų faktorinę analizę tinka gerai (KMO 0,851), tuo tarpu pasekmių faktorinę analizę tinkama pakenčiamai (KMO 0,695).

Siekiant įvertinti faktorinės analizės pagrįstumą, Bartlett'o sferiškumo testo rezultatai taip pat yra itin svarbūs. Remiantis 19-os lentelės duomenimis, abiejų pateiktų skalių Bartletto sferiškumo testo p reikšmė yra <0,0001, o tai rodo, kad faktorinę analizę šioms skalėms yra tinkama.

Analizuojant bendrumų lentelę (angl. *Communalities*), kuri pateikta 6 priede, matyti, kad tiek veiksmų, tiek pasekmių kategorijose esantys teiginiai turi daugiau nei 0,2 reikšmę „extraction“ stulpelyje. Tai rodo, kad visi teiginiai yra tinkami tolesnei faktoriaus analizei. Jei kurios nors reikšmės būtų mažesnės nei 0,2, atitinkamus teiginius reikėtų pašalinti.

Faktorių sukimui buvo naudotas „Varimax“ sukimas, kuris išlaiko faktorių nepriklausomumą (ortogonalumą), todėl faktoriai lieka nekoreliuoti. Tai ypač naudinga, kai norima aiškių, atskirų faktorių, kurie reprezentuoja skirtingas konstrukcijas ar dimensijas (Almquist et al., 2019). Šio tyrimo kontekste, atsižvelgiant į tai, jog analizuojama sąlyginai nauja tema, buvo siekiama išskirti aiškias dimensijas iš daugybės kintamųjų, todėl šis sukimo metodas laikomas tinkamu. Duomenų faktorinės analizės išskyrimo metodu pasirinkta pagrindinių ašių faktORIZACIJA (angl. *principal axis factoring*), nes ja siekiama identifikuoti nematomas (latentines) konstrukcijas, kurios gali paaiškinti stebimų kintamųjų tarpusavio ryšius, o ne vien sumažinti kintamųjų skaičių. Šis metodas taip pat tinka, kai duomenų pasiskirstymas nėra visiškai normalus (Almquist et al., 2019).

Atlikus faktorinę analizę skalėms, priskirtoms veiksmų kategorijai, pastebima itin ženklų skirtumu tarp teorijoje išskirtų konstruktų ir tai, kaip šiuos konstruktus vertino tyrimo dalyviai. Detalus teiginių pasiskirstymas pagal faktorius pateiktas 19-oje lentelėje. Prieš atliekant faktorinę analizę, teiginiai, pažymėti (r) buvo perkoduoti atvirkščiai. Pilkai pažymėti laukai atspindi faktorių, kuriam konkretus teiginys buvo priskirtas. Jei teiginiui faktorius nebuvo priskirtas, vadinasi jo faktorinis svoris nesiekė reikalingos 0,4 ribos.

19 lentelė. Veiksmų faktorinė analizė (N=461)

Kintamasis	Teiginys	Faktoriai ir jų faktoriniai svoriai						
		1	2	3	4	5	6	7
Požiūris į DI								
DI kaip bendra- autorius	Žvelgiu į DI kaip į bendraautorių, padedantį perteikti informaciją	0,112	-0,054	-0,009	0,26	0,078	0,577	0,122
	Žmogus, pasitelkdamas DI įrankius, kuria kokybiškesnius įrašus socialiniams tinklams	0,267	-0,026	0,105	0,011	-0,146	0,606	0,097
	Manau, jog DI yra lygiavertis partneris kuriant turinį	0,178	0,005	0,08	0,056	0,109	0,584	-0,166
DI kaip įrankis	DI yra įrankis, kuriuo naudojasi žmogus	0,056	0,075	0,14	0,543	-0,115	0,157	0,163
	Manau, kad žmogus turi tiesiogiai kontroliuoti, koki turini kuria DI	0,045	-0,08	-0,053	0,554	-0,125	0,051	0,017

Autentiškumo vertinimas								
Sąžiningumas	Įrašas atrodė nuoširdus	0,549	0,109	0,12	0,096	0,066	0,159	0,075
	Įrašė manipuliuojama statistiniais faktais (r)	0,07	0,438	0,001	-0,086	0,211	-0,008	0,002
	Įrašė nėra faktinių klaidų	0,318	0,023	0,23	0,227	-0,017	0,125	-0,017
Emocinis ryšys	Įrašė išvelgiu tikrosios žmogiškosios patirties	0,596	0,136	0,009	0,215	0,009	0,11	0,039
	Įrašas demonstruoja gyvybingumą ir turi savą asmenybę	0,794	-0,067	0,116	0,103	0,201	0,141	-0,088
Nuoseklumas	Įrašas pasižymi nuoseklumu	0,457	0,228	0,484	0,34	-0,111	0,031	-0,04
	Įrašo formatas ir struktūra yra nuspėjami	0,122	0,089	0,263	0,305	-0,351	-0,062	0,007
Keistumas	Įrašas yra šiurpus (r)	-0,05	0,586	-0,038	0,051	0,028	0,121	0,167
	Įrašas yra neįprastas (r)	-0,035	0,635	-0,114	0,089	-0,073	-0,117	-0,003
Sintetiškumas	Įrašas, kaip vaizdo ir teksto visuma, sudaro įspūdį, jog yra sudėliotas iš nesusijusių dedamųjų (r)	0,159	0,649	0,155	-0,136	0,13	-0,074	0,001
	Kai kurios įrašo detalės atrodo nenatūraliai (r)	0,206	0,334	0,146	-0,303	0,53	0,01	-0,138
	Įrašas pasižymi sintetiškumu (r)	0,122	0,183	-0,003	-0,144	0,744	0,049	0,027
	Įrašas, kaip visuma, sudaro nesuderintų derinių įspūdį (r)	0,111	0,575	0,216	-0,015	0,445	0,03	-0,044
Kokybės vertinimas								
Tikslumas	Įrašas atrodė faktiškai teisingas	0,197	-0,082	0,23	0,271	0,023	0,111	0,211
	Papildomai tikrinčiau pagrindinius teiginius, pateiktus įrašė (r)	-0,046	0,175	0,123	-0,354	0,308	-0,082	-0,104
Įtraukimas	Man buvo lengva išlaikyti dėmesį skaitant įrašą	0,593	0,13	0,261	0,016	0,055	-0,045	0,348
	Įrašą peržvelgiau nuo pradžios iki pabaigos	0,202	0,009	0,042	0,298	-0,082	-0,004	0,573
	Įrašas yra įdomus	0,512	0,13	0,256	-0,077	0,081	0,094	0,421
Aktualumas	Įrašas atkreipia dėmesį į kultūrinį kontekstą	0,515	-0,044	0,164	-0,027	0,021	0,115	0,087
	Įrašas orientuotas į temas, kurios aktualios dabar	0,406	0,199	0,097	0,302	-0,161	0,221	0,295
	Įrašas turi išliekamosios vertės	0,489	-0,137	0,163	-0,182	0,046	0,21	0,142
	Įrašė pastebėjau dar negirdėtų minčių	0,318	-0,259	0,134	0,071	-0,007	0,189	0,101
Suprantamumas	Įrašą buvo lengva suprasti	0,411	0,332	0,249	0,127	-0,05	-0,035	0,403
	Įrašas pateikiamas logiškai ir struktūruotai	0,505	0,254	0,335	0,097	-0,143	0,042	0,298
Pasitikėjimas								
Įrašo autoriumi galima pasitikėti		0,328	-0,015	0,64	-0,024	0,086	0,085	0,05
Įrašas yra patikimas		0,336	-0,048	0,629	-0,014	0,068	0,12	0,194

Faktorinės analizės rezultatai atskleidė, kad keturi teiginiai – susiję su įrašo nuspėjamumu, tikslumu ir negirdėtomis mintimis – nebuvo priskirti jokiui faktoriui, nes jų faktoriai svoriai nesiekė 0,4 (Almquist et al., 2019). Visi likusieji teiginiai reikšmingai prisidėjo prie vieno iš septynių identifikotų faktorių. Šie septyni faktoriai kartu paaiškina 56,58 % pradinės dispersijos (žr. 6 priedą).

Tik dviejų dimensijų – pasitikėjimo ir požiūrio į dirbtinį intelektą – teiginių pasiskirstymas atitiko iš anksto numatytą teorinę struktūrą. 6-asis faktorius atspindi požiūrį į DI kaip į bendraautorį, 4-asis – kaip į įrankį, o 3-iasis – pasitikėjimą turiniu. Tuo tarpu turinio vertinimo dimensijos išsiskyrė kitaip nei tikėtasi. Pradinėje skalėje autentiškumo dimensija buvo sudaryta iš sąžiningumo, emocinio ryšio,

nuoseklumo, keistumo ir sintetiškumo, o kokybė buvo vertinama pagal tikslumą, įtraukimą, aktualumą ir suprantamumą.

Analizės duomenimis, 5-asis faktorius atspindi sintetiškumą, nors aprėpia tik tris iš keturių numatytų teiginių. Keistumą geriausiai apibūdina 2-asis faktorius, kuris be keistumo teiginių taip pat įtraukia po vieną teiginį iš sąžiningumo ir sintetiškumo dimensijų. Vis dėlto, nuspręsta teiginį „Įrašė manipuliuojama statistiniais faktais“ pašalinti iš antrojo faktoriaus, nes teoriškai jis priklauso sąžiningumo dimensijai ir semantiškai nesusijęs su kitais keistumą atspindinčiais teiginiais, taip trikdydamas faktoriaus interpretacinį nuoseklumą. 7-asis faktorius apima vienintelį teiginį apie įrašo įtraukimą, todėl jis nebus toliau analizuojamas. Visi likę teiginiai telkiasi į pirmąjį faktorių, kuris apima sąžiningumo, emocinio ryšio, nuoseklumo, įtraukimo, aktualumo ir suprantamumo aspektus. Šių dimensijų susilieėjimas rodo, kad jos greičiausiai atstovauja bendrą, vieningą latentinį konstruktą. Dėl šios priežasties visus šiuos teiginius buvo nuspręsta sujungti į vieną skalę pavadinimu: „veržlumas ir įtaigumas“. Veržlumas apima sąžiningumą (drąsą būti savimi), emocinį ryšį (nuoširdų įtraukimą) ir nuoseklumą (pastovų tikslo siekį). Tuo tarpu įtaigumas bendrai apibūdina turinio gebėjimą įtraukti, būti aktualiam ir suprantamam.

Atnaujinus skales pagal faktorinės analizės rezultatus, išskiriamos 6 naujos skalės, priskiriamos veiksniams (žr. 20 lentelę).

20 lentelė. Atnaujintos skalės, priskiriamos veiksniams (N=461)

Kintamasis	Dimensija	Teiginių skaičius	KMO	Bartlet'o sferiškumo rodiklis
Požiūris į DI	DI kaip bendraautorius	3	0,874	< 0,001
	DI kaip įrankis	2		
Autentiškumo ir kokybės vertinimas	Veržlumas ir įtaigumas	12		
	Keistumas	3		
	Sintetiškumas	3		
Pasitikėjimas		2		

Toliau atliekama faktorinė analizė skalei, priskirtai pasekmėms. Iš 21 lentelėje pateiktų faktorinės analizės rezultatų pastebima, jog faktorių svoriai atitinka teoriškai prognozuotą struktūrą.

21 lentelė. Pasekmių faktorinė analizė

Kintamasis	Teiginys	Faktorius ir jo faktorinis svoris
		1
Elgsena	Tikėtina, jog perskaityčiau šį įrašą socialiniuose tinkluose	0,473
	Tikėtina, jog apsilankyčiau įrašo autoriaus paskyroje	0,668
	Tikėtina, jog pasidalinčiau šiuo įrašu su draugais socialiniuose tinkluose	0,714
	Tikėtina, jog sureaguočiau į šį turinį paspausdama (-s) „patinka“	0,768

Šie faktoriai paaiškina 57,33 proc. dispersijos originaliuose kintamuosiuose (žr. 6 priedą).

Atlikus tyrimo konstruktų faktorinę analizę ir atnaujinus skales, toliau nustatomas kiekvienos iš skalių patikimumams, nustatant jų KRONBACH'o alfa koeficientą per *Analyze > Scale > Reliability*

Analysis (Almquist et al., 2019). Remiantis Piligrimienės (2016) įžvalgomis, skalės laikomos patikimomis, kai jų Kronbach'o alfa koeficientas yra daugiau nei 0,6. Visgi, svarbu pabrėžti, kad patikimumas, patenkantis į intervalą tarp 0,6 ir 0,7 yra abejotino patikimumo. Pirminiai skalių patikimumo rezultatai atspindi 22-oje lentelėje.

22 lentelė. Tyrimo skalių patikimumo vertinimas (N=461)

Kintamasis	Dimensijos	Teiginių skaičius	Kronbach'o alfa koeficientas
Požiūris į DI	DI kaip bendraautorius	3	0,666
	DI kaip įrankis	2	0,611
Autentiškumo ir kokybės vertinimas	Veržlumas ir įtaigumas	12	0,869
	Keistumas	3	0,658
	Sintetiškumas	3	0,742
Pasitikėjimas		2	0,750
Elgsena		4	0,747

Skalės „Požiūris į DI“ dimensijų „DI kaip bendraautorius“ ir „DI kaip įrankis“ patikimumas, įvertintas Kronbach alfa koeficientais 0,666 ir 0,611 atitinkamai, nesiekė rekomenduojamos 0,7 ribos, todėl rezultatus reikia vertinti atsargiai. Mažesnis patikimumas gali būti siejamas su ribotu teiginių skaičiumi (po 2-3 kiekvienai dimensijai) ir dėl galimo įvairiapusio respondentų supratimo apie DI ir jo vaidmenį. Nepaisant riboto patikimumo, gauti rezultatai leidžia preliminariai apibūdinti bendrąsias tendencijas požiūryje į DI (Taber, 2018).

Vertinant autentiškumo ir kokybės matavimo skales, „Veržlumo ir įtaigumo“ dimensijos Kronbach'o alfa koeficientas siekė 0,869, rodydamas labai gerą vidinį nuoseklumą. „Sintetiškumo“ dimensijoje Kronbach'o alfa reikšmė buvo 0,742, taip pat atitinkanti priimtina patikimumo lygį. Tuo tarpu „Keistumo“ dimensijoje patikimumo rodiklis siekė tik 0,658, nesiekdamas rekomenduojamos 0,7 ribos. Žemesnės patikimumo reikšmės gali būti paaiškinamos ribotu teiginių skaičiumi, kadangi, kaip teigia Taber'is (2018), kai skales sudaro vos keli teiginiai, Kronbach'o alfa rodiklis gali būti natūraliai sumažintas.

Galiausiai, abu konstruktai – „patikimumas“ ir „elgsena“ – pasižymi priimtinais arba gerais vidinio patikimumo rezultatais, nepaisant sąlyginai mažo teiginių skaičiaus. Patikimumo skalės Kronbach'o alfa koeficientas siekia 0,750, o elgsenos – 0,747. Abi šios skalės yra patikimos.

Bendras klausimyno Kronbach'o alfa koeficientas 0,861 (teiginių skaičius – 29).

Vadovaujantis rekomendacijomis, pateiktomis metodinėje literatūroje, bei faktorinės ir patikimumo analizės rezultatais, aukščiau įvardintos matavimo skalės laikytinos tinkamomis naudoti tiek koreliacinėse, tiek regresinėse analizėse, atitinkančiomis empirinius tyrimo tikslus (Almquist et al., 2019; Briggs et al., 1980; Taber, 2018).

Dėl atliktų pakeitimų turinio autentiškumo ir kokybės vertinimo skalėse, atnaujintos H1, H2, H9 ir H10 hipotezės. Toliau pateikiamas atnaujintas hipotezių sąrašas (žr. 23 lentelę).

23 lentelė. Atnaujintų hipotezių sąrašas

Nr.	Atnaujinta hipotezė
H1	Vartotojų suvokiamas turinio autentiškumas ir kokybė teigiamai veikia jų pasitikėjimą turiniu.

	H1a	Vartotojų suvokiamas turinio veržlumas ir įtaigumas daro teigiamą įtaką turinio autentiškumo ir kokybės vertinimui.
	H1b	Vartotojų suvokiamas turinio keistumas daro neigiamą įtaką turinio autentiškumo ir kokybės vertinimui.
	H1c	Vartotojų suvokiamas turinio sintetiškumas daro neigiamą įtaką turinio autentiškumo ir kokybės vertinimui.
H2		Tikrojo turinio autoriaus atskleidimas moderuoja pasitikėjimą turiniu.
	H2a	Dirbtinio intelekto kurto turinio autoriaus atskleidimas daro neigiamą įtaką pasitikėjimui turiniu.
	H2b	Žmogaus kurto turinio autoriaus atskleidimas daro teigiamą įtaką pasitikėjimui turiniu.
H3		Vartotojų pasitikėjimo turiniu pokytis teigiamai veikia jų polinkį sąveikauti su turiniu.
H4		Požiūris į DI moderuoja kaip vartotojų pasitikėjimo pokytį veikia tikrojo turinio autoriaus atskleidimas.
	H4a	Dirbtinio intelekto turinio tikrojo autoriaus atskleidimas neigiamai veikia vartotojų, kurie į DI žvelgia kaip į bendraautorį, pasitikėjimo turiniu pokyčiui.
	H4b	Dirbtinio intelekto turinio tikrojo autoriaus atskleidimas teigiamai veikia vartotojų, kurie į DI žvelgia kaip į įrankį, pasitikėjimo turiniu pokyčiui.
H5		Vartotojų žinios apie DI įrankius teigiamai veikia jų gebėjimą identifikuoti tikrąjį turinio autorių.
H6		Vartotojų suvokiamas gebėjimas tinkamai identifikuoti DI turinį teigiamai veikia jų gebėjimą identifikuoti DI turinį.
H7		Vartotojų suvokiamas pasitikėjimo turiniu didėjimas žinant autorių teigiamai veikia jų pasitikėjimą turiniu, atskleidus tikrąjį turinio autorių.
H8		Žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio autentiškumui ir kokybei nei dirbtinio intelekto sukurtas turinys.
	H8a	Žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio veržlumui ir įtaigumui nei dirbtinio intelekto sukurtas turinys.
	H8b	Žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnę neigiamą poveikį vartotojų suvokiamam turinio keistumui nei dirbtinio intelekto sukurtas turinys.
	H8c	Žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnę neigiamą poveikį vartotojų suvokiamam turinio sintetiškumui nei dirbtinio intelekto sukurtas turinys.
H9		Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų pasitikėjimui turiniu nei dirbtinio intelekto sukurtas turinys.

Interpretuojant faktorinės analizės rezultatus, galima pastebėti, jog trečiame skyriuje parengta apklausos anketa ir skirtingus konstruktus išmatuojančios skalės yra patikimos. Visgi, svarbu pabrėžti, kad po faktorinės analizės buvo atskiros skalės, skirtos matuoti turinio autentiškumui ir kokybei buvo itin reikšmingai modifikuotos. Vietoje 9 dimensijų (sąžiningumas, emocinis ryšys, nuoseklumas, keistumas, sintetiškumas, tikslumas, įtraukimas, aktualumas, suprantamumas), sudarančių šias skalės, tolimesnėje analizėje bus naudojamos trys: veržlumas ir įtaigumas, keistumas, sintetiškumas. Likusios skalės po faktorinės analizės išlaikė teorinę struktūrą.

4.3. Psichometriniais matavimais vertinamų kintamųjų charakteristikos

Prieš pereinant prie iškeltų hipotezių tikrinimo, svarbu atlikti psichometriniais matavimais vertinamų kintamųjų charakteristikų aprašomąją analizę. Šiuo atveju išskiriami aktualių matavimo skalių minimali bei maksimali reikšmė, vidurkis ir standartinis nuokrypis. Prieš analizę, kiekviena skalė, kuri buvo sudaryta iš daugiau nei vieno teiginio, buvo suvidurkinta, pasitelkiant suminio vidurkio metodą (Piligrimienė, 2016). Bendros kintamųjų charakteristikos atsispindi 24-oje lentelėje.

24 lentelė. Psichometriniais matavimais vertinamų kintamųjų charakteristikos (N=461)

Kintamasis	Dimensijos (teiginių skaičius)	Minimali reikšmė	Maksimali reikšmė	Vidurkis	Standartinis nuokrypis
Požiūris į DI	DI kaip bendraautorius (3)	1	5	3,15	0,76
	DI kaip įrankis (2)	1	5	4,16	0,78
Turinio autentiškumas ir kokybė	Įtaigumas ir veržlumas (12)	1	4,83	3,1	0,61
	Keistumas (3)	1	5	2,6	0,78
	Sintetiškumas (3)	1	5	3,11	0,74
Pasitikėjimas (2)		1	5	2,90	0,73
Pasitikėjimo pokytis (1)		1	5	3,13	1,04
Elgsena (4)		1	5	2,51	0,82
Žinios apie DI įrankius (1)		1	5	2,6	1,19
Suvokiamas gebėjimas atpažinti DI turinį (1)		1	5	3,21	1,09
Suvokiamas pasitikėjimas žinant autorių (1)		1	5	4,02	0,96

Remiantis lentelėje pateiktais duomenimis, pastebima, jog DI kaip bendraatoriaus kintamojo vidurkis yra 3,15. Pagal matavimo skalę, ši reikšmė artima „nei sutinku, nei nesutinku“ (3). Tuo tarpu DI kaip įrankio kintamojo vidurkis siekia 4,16 – ši reikšmė artima „sutinku“ (4) pasirinkimo variantui. Sąlyginai mažas standartinis nuokrypis (0,76 ir 0,78 atitinkamai) vertinant abi šias dimensijas indikuoja, jog respondentų atsakymai buvo gana nuoseklūs. Visgi, pagal šių dviejų dimensijų vidurkių skirtumus matyti, kad respondentai labiau linkę sutikti, jog DI yra įrankis, nei manyti, kad DI yra bendraautorius. Tai gali reikšti, jog vartotojai socialiniuose tinkluose linkę DI matyti kaip pagalbinę priemonę, o ne lygiavertį kūrėją.

Turinio autentiškumo ir kokybės skalėje įtaigumo ir veržlumo kintamojo vertinimo vidurkis yra 3,1, o standartinis nuokrypis – 0,61. Šiuo atveju vidurkis artimas „nei sutinku, nei nesutinku“ (3) pasirinkimo variantui. Panašus vidurkis pastebimas ir žvelgiant į sintetiškumo kintamojo vertinimą – čia jis siekia 3,11, o tai taip pat indikuoja, jog respondentų vertinimas daugiausia buvo neutralus, atitinkantis „nei sutinku, nei nesutinku“ opciją. Tuo tarpu įrašo keistumo vidurkis yra 2,6 – iš dalies artimas tiek „nesutinku“ (2), tiek „nei sutinku, nei nesutinku“ (3) opcijai. Standartinis nuokrypis visais trimis atvejais gan mažas (0,61, 0,74 ir 0,78 atitinkamai) – dėl to galima manyti, jog variacija tarp atsakymų nebuvo didelė.

Pasitikėjimo kintamojo vidurkis yra 2,90. Pagal matavimo skalę ši reikšmė artima atsakymui „nei sutinku, nei nesutinku“ (3), tačiau yra šiek tiek žemiau šios neutralios ribos, todėl galima daryti prielaidą, kad respondentai linkę šiek tiek nesutikti su teiginiais, apibūdinančiais pasitikėjimą turiniu. Standartinis nuokrypis yra 0,73, o tai rodo santykinai mažą atsakymų variaciją. Tuo tarpu pasitikėjimo pokyčio vidurkis yra 3,13, o pagal matavimo skalę ši reikšmė artima „pasitikėjimas turiniu nepasikeitė“ (3) atsakymo variantui, su tendencija į didesnio pasitikėjimo kryptį. Standartinis nuokrypis siekia 1,04 – tai sąlyginai aukšta reikšmė, kuri, tikėtina, gali būti pagrįsta eksperimentiniu tyrimo metodu.

Elgsenos kintamojo vidurkis yra 2,51. Pagal matavimo skalę ši reikšmė yra tarp atsakymų „nesutinku“ (2) ir „nei sutinku, nei nesutinku“ (3), tačiau artimesnė „nesutinku“ pasirinkimui. Tai rodo, kad respondentai yra labiau linkę nesąveikauti su turiniu. Standartinis nuokrypis siekia 0,82, o tai reiškia, jog egzistavo tam tikras respondentų nuomonių išsiskyrimas, tačiau jis nebuvo ekstremalus.

Respondentų suvokiamo gebėjimo atpažinti DI turinį vidurkis yra 3,21, o tai rodo artimą neutralumui vertinimą – respondentai nei stipriai pasitiki, nei nepasitiki savo gebėjimu atpažinti DI generuotą turinį, o gana aukštas standartinis nuokrypis (1,09) rodo respondentų nuomonių skirtumus. Tuo tarpu, pasitikėjimo turiniu, žinant jo autorių, kintamojo vidurkis (4,02) taip pat artimas „sutinku“ (4) atsakymui, o tai rodo tendenciją didesniai pasitikėjimui, kai autorius yra žinomas. Tuo tarpu standartinis nuokrypis (0,96) leidžia spręsti apie išsiskiriančią respondentų vertinimą šiuo klausimu.

Toliau siekiama nustatyti, kokią įtaką respondentų sociodemografinės charakteristikos daro jų vertinimui tyrimo skalių kontekste. Pirmiausia, siekiant įvertinti, ar lytis turi reikšmingos įtakos skalių vertinimui, buvo atliktas Mann–Whitney U neparametrinis statistinis testas, skirtas dviejų nepriklausomų grupių pasiskirstymų palyginimui, kai duomenys nėra normaliai pasiskirstę, o tiriamasis kintamasis yra kategorinis ar skalės su nenormaliu skirstiniu (Taber, 2018). Dėl itin mažos kitos lyties atstovų imties ($N = 1$), šios grupės duomenys laikomi nereprezentatyviais, todėl pagrindinė analizė orientuota į vyrų ir moterų palyginimą. Testo rezultatai pateikiami 25 lentelėje.

25 lentelė. Kintamųjų reikšmių priklausomybė nuo lyties ($N=460$)

Kintamasis	Dimensijos	Lytis	Vidutinis rangas	p-reikšmė
Požiūris į DI	DI kaip bendraautorius	Vyras	226,45	0,556
		Moteris	233,73	
	DI kaip įrankis	Vyras	228	0,711
		Moteris	232,49	
Turinio autentiškumas ir kokybė	Įtaigumas ir veržlumas	Vyras	234,12	0,602
		Moteris	227,62	
	Keistumas	Vyras	217,09	0,051
		Moteris	241,19	
	Sintetiškumas	Vyras	225,11	0,432
		Moteris	234,8	
Pasitikėjimas		Vyras	225,67	0,464
		Moteris	234,35	
Pasitikėjimo pokytis		Vyras	220,10	0,114
		Moteris	238,79	
Elgsena		Vyras	208,04	0,001
		Moteris	248,4	
Žinios apie DI įrankius		Vyras	224,76	0,394
		Moteris	235,07	
Suvokiamas gebėjimas atpažinti DI turinį		Vyras	243,97	0,044
		Moteris	219,77	

Suvokiamas pasitikėjimas žinant autorių	Vyras	233,17	0,681
	Moteris	228,37	

Iš lentelėje pateiktų rezultatų galima pastebėti, jog lytis turi statistiškai reikšmingą poveikį vartotojų elgsenai ir suvokiamam dirbtinio intelekto pagalba kurto turinio objektyvumui, nes čia p reikšmė yra mažesnė nei 0,05. Remiantis Mann–Whitney U testo rezultatais, moterys (vidutinis rangas 248,4) yra labiau linkusios sąveikauti su turiniu nei vyrai (vidutinis rangas 208,04). Tuo tarpu vyrai (vidutinis rangas 243,97) labiau nei moterys (vidutinis rangas 219,77) mano gebantys atpažinti, ar turinys buvo sukurtas dirbtinio intelekto pagalba.

Žvelgiant į rezultatus taip pat pastebėta, jog kintamųjų, matuojančių požiūrį į DI, turinio autentiškumą ir kokybę, respondentų pasitikėjimą, žinias apie DI įrankius, suvokiamą DI objektyvumą, DI turinio žymėjimo poreikį bei suvokiamą pasitikėjimą žinant turinio autorių, vertinimas nepriklauso nuo lyties.

Kadangi amžius, išsilavinimas, socialinių tinklų naudojimo dažnis ir trukmė bei DI įrankių naudojimo dažnis yra tvarkos skalės kintamieji, jiems – priešingai nei lyčiai – svarbi reikšmių išsidėstymo tvarka. Dėl to, siekiant nustatyti jų tarpusavio ryšį, atliekama koreliacijos analizė. Norint pasirinkti tinkamą koreliacijos koeficientą – Pearson'o ar Spearman'o – būtina patikrinti, ar kintamųjų reikšmės pasiskirsčiusios pagal normalųjį skirstinį. Tam naudojamas Kolmogorovo-Smirnov (K-S) testas. Remiantis analizės rezultatais (žr. 26 lentelę), visų tirtų kintamųjų p reikšmės yra mažesnės nei 0,05. Tai rodo, kad jų skirstiniai statistiškai reikšmingai skiriasi nuo normaliojo skirstinio.

26 lentelė. Kolmogorov'o-Smirnov'o (K-S) normalumo testo rezultatai (N=461)

Kintamasasis	p-reikšmė
DI kaip bendraautorius	<0,001
DI kaip įrankis	<0,001
Įtaigumas ir veržlumas	<0,001
Keistumas	<0,001
Sintetiškumas	<0,001
Pasitikėjimas	<0,001
Pasitikėjimo pokytis	<0,001
Elgsena	<0,001
Žinios apie DI įrankius	<0,001
Suvokiamas gebėjimas atpažinti DI turinį	<0,001
Suvokiamas pasitikėjimas žinant autorių	<0,001
Amžius	<0,001
Išsilavinimas	<0,001
Socialinių tinklų naudojimo dažnis	<0,001
Socialinių tinklų naudojimo trukmė	<0,001
DI įrankių naudojimo dažnis	<0,001

Atsižvelgiant į tai, jog kintamieji nėra pasiskirstę pagal normalųjį skirstinį, koreliacinei analizei buvo naudojamas Spearman'o koreliacijos koeficientas. Svarbu paminėti, kad socialinių tinklų ir DI įrankių naudojimo dažnis buvo užkoduoti nuo 1 iki 7, kur 1 atitinka didžiausią naudojimo dažnį, o 7

– mažiausią; tuo tarpu socialinių tinklų naudojimo trukmė užkoduota nuo 1 iki 4, kur 1 žymima mažiausia trukmė, o 4 – didžiausia (žr. 2 priedą). Koreliacinės analizės rezultatai pateikiami 27-oje lentelėje.

27 lentelė. Kintamųjų reikšmių priklausomybė nuo amžiaus, išsilavinimo, socialinių tinklų naudojimo dažnio, trukmės bei DI įrankių naudojimo dažnio (N=461)

Skalė	Dimensijos	Aprašomasis kintamasis				
		Amžius	Išsilavinimas	Soc. dažnis	Soc. trukmė	DI dažnis
Požiūris į DI	DI kaip bendraautorius	0,158**	0,058	-0,023	0,01	-0,079
	DI kaip įrankis	0,019	0,036	-0,004	-0,004	-0,086
Turinio autentiškumas ir kokybė	Įtaigumas ir veržlumas	0,047	-0,041	-0,01	-0,005	-0,097*
	Keistumas	0,083	-0,054	-0,019	-0,003	0,101*
	Sintetiškumas	0,026	0,081	0,003	-0,011	0,024
Pasitikėjimas		0,06	-0,091*	0,096*	0,08	-0,05
Pasitikėjimo pokytis		0,015	0,027	-0,059	0,041	0,044
Elgsena		0,161**	0,017	-0,041	0,021	0,042
Žinios apie DI įrankius		-0,205**	0,009	0,06	0,100*	-0,250**
Suvokiamas gebėjimas atpažinti DI turinį		-0,210**	0,057	-0,132**	-0,014	-0,07
Suvokiamas pasitikėjimas žinant autorių		0,096*	0,035	-0,031	-0,016	0,04

*p<0,05, **p<0,01

Iš koreliacinės analizės rezultatų, pateiktų 27-oje lentelėje, matoma, jog yra statistiškai reikšmingas teigiamas labai silpnas ryšys tarp požiūrio į DI kaip į bendraautorius ir amžiaus ($r = 0,158$). Vadinasi, didėjant respondentų amžiui, jų polinkis žvelgti į DI kaip į bendraautorius taip pat didėja. Taip pat pastebima statistiškai reikšminga labai silpna teigiama koreliacija tarp amžiaus ir ketinimo sąveikauti su turiniu ($r = 0,161$) – vyresni respondentai yra labiau linkę komentuoti, apsilankyti autoriaus paskyroje, perskaityti ir sureaguoti į turinį nei jaunesni. Kitą vertus, tarp amžiaus ir žinių apie DI įrankius pastebimas statistiškai reikšmingas silpnas neigiamas ryšys ($r = 0,205$), indikuojantis, jog jaunesni respondentai žino daugiau apie DI įrankius nei vyresni. Statistiškai reikšmingas neigiamas labai silpnas ryšys pastebimas ir tarp amžiaus bei suvokiamo gebėjimo atpažinti DI turinį nuo žmogaus ($r = -0,210$) – tai leidžia daryti prielaidą, jog jaunesni asmenys mano, jog DI kurtą turinį gali atskirti lengviau nei vyresnieji. Suvokiamas pasitikėjimas žinant turinio autorių priklauso nuo amžiaus – tarp šių kintamųjų pastebėtas statistiškai reikšmingas labai silpnas teigiamas ryšys ($r = 0,096$), imponuojantis, jog vyresni respondentai yra labiau linkę pasitikėti turiniu, jei žino, kas yra jo autorius.

Tarp išsilavinimo ir pasitikėjimo turiniu pastebėtas statistiškai reikšmingas labai silpnas neigiamas ryšys ($r = -0,091$). Tai reiškia, jog aukštesnio išsilavinimo asmenys yra linkę mažiau pasitikėti turiniu socialiniuose tinkluose.

Tarp socialinių tinklų naudojimo dažnio ir pasitikėjimo pastebėtas silpnas statistiškai reikšmingas ryšys ($r = 0,096$), kuris imponuoja, kad respondentai, kurie socialiniais tinklais naudojami rečiau, yra linkę labiau pasitikėti ten esančiu turiniu. Tuo tarpu silpnas neigiamas ryšys socialinių tinklų naudojimo dažnio ir suvokiamo gebėjimo atpažinti DI turinį ($r = -0,132$), imponuoja jog asmenys, kurie dažniau lankosi socialiniuose tinkluose yra labiau linkę manyti, jog gali lengvai atpažinti DI pagalba sukurtą turinį.

Asmenys, daugiau laiko praleidžiantys socialiniuose tinkluose, apie DI įrankius žino daugiau, nei asmenys, kurie tai daro rečiau. Šią prielaidą galima daryti atkreipus dėmesį į statistiškai reikšmingą silpną ryšį tarp žinių apie DI įrankius ir socialinių tinklų naudojimo trukmės ($r = 0,100$).

Silpnas neigiamas ryšys tarp dirbtinio intelekto įrankių naudojimo dažnumo bei veržlumo ir įtaigumo ($r = -0,097$) suponuoja, kad respondentai, kurie dirbtinio intelekto įrankius naudoja dažniau, yra linkę turinį priimti kaip įtaigesnį ir veržlesnį. Silpnas teigiamas ryšys tarp turinio keistumo ir DI įrankių naudojimo dažnio ($r = 0,101$) suponuoja, kad respondentai, kurie DI įrankius naudoja rečiau, yra linkę pateiktą turinį suvokti kaip keistesnį. Galiausiai, silpnas neigiamas ryšys ($r = -0,250$) suponuoja, kad respondentai, kurie DI įrankius naudoja dažniau, taip pat apie juos daugiau žino.

Kitų reikšmingų ryšių tarp kintamųjų nebuvo rasta.

4.4. Kintamųjų tarpusavio ryšių analizė ir hipotezių tikrinimas

Šioje baigiamojo magistro darbo dalyje tiriami pagrindinių kintamųjų ryšiai, tikrinamos pagrindinės tyrimo hipotezės. Pasitelkiant koreliacinę analizę, siekiama išskirti, kokie ryšiai tarp tyrimui aktualių kintamųjų egzistuoja. Pagal 27 lentelėje ir 8 priede pateiktus normalumo testo rezultatus nustatyta, jog nei vienas kintamasis nėra pasiskirstęs pagal normalųjį skirstinį, todėl koreliacinei analizei naudojamas Spearman'o koreliacijos koeficientas. Žemiau pateiktoje lentelėje (žr. 28 lentelę) išskiriami kintamieji bei jiems priskiriamas eilės numeris, padėsiantis lengviau interpretuoti koreliacinių ryšių lentelę.

28 lentelė. Kintamųjų kodavimas koreliacinei analizei

Nr.	Kintamieji
1	DI kaip bendraautorius
2	DI kaip įrankis
3	Įtaigumas ir veržlumas
4	Keistumas
5	Sintetiškumas
6	Pasitikėjimas
7	Pasitikėjimo pokytis
8	Elgsena
9	Žinios apie DI įrankius
10	Suvokiamas gebėjimas atpažinti DI turinį
11	Suvokiamas pasitikėjimas žinant autorių

Koreliacijos stiprumui nustatyti pasitelkiama Kafle (2023) skalė (žr. 29 lentelę). Svarbu pabrėžti, kad koreliacija tarp kintamųjų nėra įrodymas, jog kintamieji yra priežastiniai, ji padeda atskleisti tik ryšio

stiprumą ir kryptį. Vidutinė ar labai silpna koreliacija gali būti laikoma priimtina, jei tyrimo tema yra sudėtinga (Kafle, 2023).

29 lentelė. Koreliacijos koeficientų interpretavimas pagal Kafle (2023)

Koreliacijos koeficientas, r		
$0 < r \leq 0,4$	$0,4 < r \leq 0,7$	$0,7 < r \leq 1$
Silpna koreliacija	Vidutinė koreliacija	Stipri koreliacija

Toliau pateikiami tyrimui svarbių kintamųjų koreliacinės analizės rezultatai (žr. 30 lentelę). Spalvotai pažymėti laukai atspindi statistiškai reikšmingas koreliacijas.

30 lentelė. Koreliacinės analizės tarp tyrimui svarbių kintamųjų rezultatai (N=461)

	1	2	3	4	5	6	7	8	9	10	11
1											
2	0,045										
3	0,303**	0,107*									
4	0,026	-0,055	-0,220**								
5	0,001	0,136**	-0,317**	0,555**							
6	0,182**	-0,057	0,509**	-0,036	-0,198**						
7	-0,015	-0,115*	-0,075	0,032	0,043	-0,05					
8	0,211**	-0,029	0,421**	0,08	-0,226**	0,246**	0,188**				
9	-0,01	0,041	-0,078	-0,056	-0,056	-0,079	0,044	-0,035			
10	0,066	0,111*	-0,016	-0,102*	0,048	-0,028	0,032	-0,075	0,116*		
11	0,256**	-0,026	0,076	0,155**	-0,091	0,055	-0,067	-0,044	-0,011	-0,056	

* $p < 0,05$, ** $p < 0,01$

Žvelgiant į koreliacinės analizės rezultatus, pastebima keletas reikšmingų sąsajų tarp kintamųjų. Egzistuoja teigiamas, silpnas ryšys tarp teksto įtaigumo ir veržlumo bei požiūrio į DI kaip į bendraautorių ($r = 0,303$, $p < 0,01$). Tai leidžia daryti prielaidą, kad žmonės, pateiktą turinį suvokę kaip įtaigų ir veržlų, yra linkę į DI žvelgti palankiau – kaip į bendraautorių. Panaši, nors ir silpnesnė tendencija stebima ir tarp įtaigumo bei veržlumo ir požiūrio į DI kaip į įrankį ($r = 0,107$), leidžianti manyti, kad suvokiamas įtaigumas ir veržlumas taip pat skatina DI suvokimą kaip įrankį, tačiau šis ryšys yra mažiau išreikštas.

Pastebėtas neigiamas, silpnas ryšys tarp įtaigumo ir veržlumo bei turinio keistumo ($r = -0,220$), rodantis, jog įtaigiai ir veržliai pateiktas turinys rečiau suvokiamas kaip keistas. Be to, reikšmingas silpnas teigiamas ryšys tarp turinio sintetiškumo ir požiūrio į DI kaip į įrankį ($r = 0,136$) leidžia suprasti, kad respondentai, turinį vertinantys kaip sintetišką, linkę DI laikyti labiau įrankiu nei kūrybiniu partneriu.

Silpnas, neigiamas ir statistiškai reikšmingas ryšys tarp įtaigumo ir veržlumo bei sintetiškumo ($r = -0,317$, $p < 0,01$) leidžia daryti prielaidą, kad kuo turinys įtaigesnis ir veržlesnis, tuo jis rečiau

vertinamas kaip sintetiškas. Teigiamas, silpnas ryšys tarp požiūrio į DI kaip į bendraautorijų ir pasitikėjimo turiniu ($r = 0,182$) rodo, kad žmonės, linkę laikyti DI bendraautoriumi, labiau pasitiki ir turiniu. Tuo tarpu tarp pasitikėjimo turiniu ir jo sintetiškumo pastebimas neigiamas, silpnas ryšys ($r = -0,198$), suponuojantis, kad kuo turinys atrodo sintetiškesnis, tuo mažesnis pasitikėjimas juo.

Toliau matomas teigiamas, silpnas ir statistiškai reikšmingas ryšys tarp požiūrio į DI kaip į bendraautorijų ir elgsenos ($r = 0,211$), reiškiantis, kad tokį požiūrį turintys respondentai yra labiau linkę sąveikauti su DI generuotu turiniu. Tuo tarpu tarp elgsenos ir turinio sintetiškumo ryšys yra neigiamas ir statistiškai reikšmingas ($r = -0,226$), leidžiantis daryti išvadą, kad kuo turinys atrodo sintetiškesnis, tuo mažesnis noras su juo sąveikauti.

Pastebėtas teigiamas, silpnas ryšys tarp pasitikėjimo turiniu ir elgsenos ($r = 0,246$), kuris parodo, kad didėjant pasitikėjimui turiniu, auga ir polinkis į aktyvią sąveiką su juo. Silpnas, nors ir labai nežymus, ryšys stebimas tarp požiūrio į DI kaip į įrankį ir suvokiamo gebėjimo atpažinti DI turinį ($r = 0,111$), o tai rodo, kad žmonės, manantys galintys atpažinti DI turinį, DI suvokia labiau kaip įrankį nei kūreją. Taip pat pastebimas neigiamas silpnas ryšys tarp pasitikėjimo pokyčio ir keistumo ($r = -0,115$), leidžiantis daryti prielaidą, kad respondentai yra mažiau linkę pasitikėti DI turiniu po tikrojo autoriaus atskleidimo, jei jie į DI žvelgia kaip į įrankį. Kitą vertus, tarp pasitikėjimo pokyčio ir elgsenos pastebima statistiškai reikšminga silpna tendencija ($r = 0,188$), suponuojanti, jog respondentai, po turinio autoriaus atskleidimo turiniu pasitikėję labiau taip pat yra labiau linkę su juo sąveikauti.

Neigiamas, silpnas ryšys ($r = -0,102$) pastebimas tarp keistumo ir gebėjimo atpažinti DI turinį, kas leidžia manyti, kad tie, kurie geriau atpažįsta DI generuotą turinį, jį vertina kaip mažiau keistą. Taip pat pastebimas teigiamas ryšys tarp žinių apie DI ir gebėjimo atpažinti DI turinį ($r = 0,116$), rodantis, kad didesnis žinių lygis padeda lengviau identifikuoti dirbtinio intelekto kurtą turinį.

Svarbu paminėti ir ryšį tarp pasitikėjimo turiniu žinant autoriaus tapatybę ir požiūrio į DI kaip į bendraautorijų ($r = 0,256$), kuris rodo, kad tie, kurie linkę pasitikėti turiniu žinodami jo autorių, taip pat labiau palankiai vertina DI kaip kūrybinį partnerį. Be to, statistiškai reikšmingas ryšys pastebėtas tarp įtaigumo ir veržlumo bei pasitikėjimo žinant autorių ($r = 0,155$), kas rodo, kad tokį turinį žmonės yra linkę laikyti patikimesniu, jei žino, kas jį sukūrė.

Vidutinio stiprumo teigiamas ryšys tarp turinio keistumo ir sintetiškumo ($r = 0,555$, $p < 0,01$) rodo, kad kuo labiau turinys atrodo sintetiškas, tuo didesnė tikimybė, kad jis bus vertinamas kaip keistas. Kitas vidutinio stiprumo teigiamas ryšys stebimas tarp įtaigumo ir veržlumo bei pasitikėjimo turiniu ($r = 0,509$, $p < 0,01$), o tai parodo, kad emocinis, nuoseklus ir įtraukiantis turinys labiau skatina auditorijos pasitikėjimą. Galiausiai, ryšys tarp įtaigumo ir veržlumo bei elgsenos ($r = 0,421$) leidžia daryti prielaidą, kad įtaigus ir veržlus turinys skatina aktyvesnį vartotojų įsitraukimą. Kitų statistiškai reikšmingų ryšių tarp kintamųjų nebuvo nustatyta.

Toliau atliekamas hipotezių tikrinimas pasitelkiant įvairius statistinius metodus: regresinė analizė (angl. *Regression Analysis*), moderavimo analizė (angl. *Moderation Analysis*), Mann-Whitney U testas, chi-kvadrato testas.

Siekiant išsiaiškinti kokią įtaką respondentų pasitikėjimui turiniu daro jo autentiškumas ir kokybė, atliekama daugialypė regresinė analizė. Modelio tinkamumas grindžiamas 31 lentelėje ir 9 priede pateiktais duomenimis. Tikrinamos šios hipotezės:

H1a. Vartotojų suvokiamas turinio veržlumas ir įtaigumas daro teigiamą įtaką turinio autentiškumo ir kokybės vertinimui.

H1b. Vartotojų suvokiamas turinio keistumas daro neigiamą įtaką turinio autentiškumo ir kokybės vertinimui.

H1c. Vartotojų suvokiamas turinio sintetiškumas daro neigiamą įtaką turinio autentiškumo ir kokybės vertinimui

31 lentelė. Daugialypė regresija tarp autentiškumo ir kokybės bei pasitikėjimo. Modelio tinkamumas (N=461)

Rodiklis	Reikšmė
R	0,595
R ²	0,354
F reikšmė	83,334
p reikšmė	< 0,001
VIF reikšmės	1,07 – 1,23

Regresijos modelis yra statistiškai reikšmingas ($F = 83,334$, $p < 0,001$) o determinacijos koeficientas ($R^2 = 0,354$) rodo, kad nepriklausomi kintamieji paaiškina apie 35,4 proc. pasitikėjimo turiniu dispersijos. Tai leidžia daryti prielaidą, kad pasirinkti kintamieji yra gana svarbūs aiškinant pasitikėjimą turiniu. Visų kintamųjų VIF reikšmės buvo mažesnės nei 1,3, o tolerancijos reikšmės viršijo 0,8, todėl galime teigti, kad tarp nepriklausomų kintamųjų nėra reikšmingo daugiakolinearumo, galinčio iškreipti regresijos koeficientų interpretaciją. 32-oje lentelėje apibūdinama nepriklausomų kintamųjų įtaka.

32 lentelė. Daugialypė regresija tarp autentiškumo ir kokybės bei pasitikėjimo. Nepriklausomų kintamųjų įtaka (N=461)

Kintamasis	β	B	t reikšmė	p reikšmė	VIF	Reikšmingumas, kryptis
Veržlumas ir įtaigumas	0,570	0,680	14,667	< 0,001	1,066	Reikšminga, teigiama
Keistumas	0,152	0,141	3,694	< 0,001	1,203	Reikšminga, teigiama
Sintetiškumas	-0,131	-0,125	-3,139	0,002	1,24	Reikšminga, neigiama

Veržlumo ir įtaigumo vertinimas turi stiprią teigiamą įtaką pasitikėjimui ($\beta = 0,570$, $p < 0,001$). Turinio keistumas turi statistiškai reikšmingą teigiamą poveikį pasitikėjimui ($\beta = 0,152$, $p < 0,001$). Didėjant turinio keistumui, didėja ir pasitikėjimas juo. Tuo tarpu turinio sintetiškumas turėjo neigiamą ir statistiškai reikšmingą poveikį pasitikėjimui ($\beta = -0,131$, $p = 0,002$).

Remiantis atliktos regresinės analizės rezultatais, hipotezės **H1a** (Vartotojų suvokiamas turinio veržlumas ir įtaigumas daro teigiamą įtaką turinio autentiškumo ir kokybės vertinimui) bei **H1c** (Vartotojų suvokiamas turinio sintetiškumas daro neigiamą įtaką turinio autentiškumo ir kokybės vertinimui) yra **patvirtinamos**. Tuo tarpu **H1b** (vartotojų suvokiamas turinio keistumas daro neigiamą įtaką turinio autentiškumo ir kokybės vertinimui) hipotezė yra **atmetama**.

Toliau siekiant išsiaiškinti kaip tikrojo turinio autoriaus atskleidimas moderuoja pasitikėjimą turiniu buvo atliktas Mann–Whitney U testas, atsižvelgiant į tai, jog tyrime dalyvavo dvi eksperimentės

grupės, o analizuojamų kintamųjų duomenys neatitinka normaliojo skirstinio. Testo rezultatai pateikiami 33-ioje lentelėje ir 10 priede. Tikrinamos šios hipotezės:

H2a. Dirbtinio intelekto kurto turinio autoriaus atskleidimas daro neigiamą įtaką pasitikėjimui turiniu

H2b. Žmogaus kurto turinio autoriaus atskleidimas daro teigiamą įtaką pasitikėjimui turiniu

33 lentelė. Neparametrinis Mann-Whitney U testas tarp turinio autoriaus atskleidimo ir pasitikėjimo pokyčio (N=461)

Rodiklis	Reikšmė
N (autorius – DI / žmogus)	232 / 229
Vidutinis rangas (autorius – DI)	171,51
Vidutinis rangas (autorius – žmogus)	291,27
Z reikšmė	-10,179
p reikšmė	< 0,001

Atliekant Mann–Whitney U testą buvo siekiama įvertinti, ar tikrojo turinio autoriaus atskleidimas – dirbtinio intelekto ar žmogaus – daro įtaką pasitikėjimo turiniu vertinimui. Testo rezultatai parodė statistiškai reikšmingą skirtumą tarp dviejų grupių ($U = 12761,5$, $Z = -10,18$, $p < 0,001$). Vidutinis rangas DI autoriaus atveju buvo 171,51, o žmogaus autoriaus – 291,27, kas rodo, jog respondentai pasitikėjo turiniu labiau tuomet, kai buvo atskleista, kad jį sukūrė žmogus, o ne dirbtinis intelektas.

Remiantis atlikto Mann-Whitney U testo rezultatais, abi hipotezės – **H2a** (dirbtinio intelekto kurto turinio autoriaus atskleidimas daro neigiamą įtaką pasitikėjimui turiniu) ir **H2b** (žmogaus kurto turinio autoriaus atskleidimas daro teigiamą įtaką pasitikėjimui turiniu) – yra **patvirtinamos**.

Toliau analizuojama kokią įtaką vartotojų pasitikėjimas turiu daro jų polinkiui su juo sąveikauti pasitelkiant paprastąją regresinę analizę. Modelio tinkamumas grindžiamas 34 lentelėje ir 11 priede pateiktais duomenimis. Tikrinama ši hipotezė:

H3. Vartotojų pasitikėjimo turiniu pokytis teigiamai veikia jų polinkį sąveikauti su turiniu.

34 lentelė. Paprastoji regresija tarp pasitikėjimo pokyčio ir elgsenos. Modelio tinkamumas (N=461)

Rodiklis	Reikšmė
R	0,207
R ²	0,043
F reikšmė	20,553
p reikšmė	< 0,001

Remiantis modelio tinkamumo vertinimo duomenimis, kad regresijos modelis yra statistiškai reikšmingas ($F = 20,553$, $p < 0,001$), o tai reiškia, kad pasitikėjimo lygis turi reikšmingą įtaką vartotojų elgsenai. Koreliacijos koeficientas tarp kintamųjų buvo $R = 0,207$, o tai rodo silpną, bet teigiamą ryšį tarp pasitikėjimo ir sąveikos. Determinacijos koeficientas $R^2 = 0,043$ rodo, kad pasitikėjimas turiniu paaiškina 4,3 % vartotojų sąveikavimo elgsenos dispersijos. Labai svarbu paminėti, kad R^2 reikšmė šiuo atveju yra itin maža, o tai gali reikšti, kad į modelį neįtraukti svarbūs

kintamieji, turintys didesnę įtaką – pasitikėjimas gali būti tik vienas iš daugelio veiksnių, darantys įtaką elgsenai. 35-oje lentelėje apibūdinama nepriklausomo kintamojo įtaka.

35 lentelė. Paprastoji regresija tarp pasitikėjimo pokyčio ir elgsenos. Nepriklausomojo kintamojo įtaka (N=461)

Kintamasis	β	B	t reikšmė	p reikšmė	Reikšmingumas, kryptis
Elgsena	0,207	0,164	4,534	< 0,001	Reikšminga, teigiama

Pasitikėjimas turiniu turi silpną teigiamą įtaką elgsenai ($\beta = 0,207$, $p < 0,001$).

Remiantis paprastosios regresijos rezultatais, hipotezė **H3** (vartotojų pasitikėjimas turiniu teigiamai veikia jų polinkį sąveikauti su turiniu) yra **sąlyginai patvirtinama** – nors regresijos modelis yra tinkamas duomenų dispersijai aiškinti, itin žema R^2 reikšmė leidžia daryti prielaidą, kad šis modelis paaiškina pernelyg mažą duomenų dalį.

Toliau atliekama moderacijos analizė, siekiant išsiaiškinti, ar ir kaip požiūris į dirbtinį intelektą keičia tikrojo turinio autoriaus atskleidimo poveikį vartotojų pasitikėjimui turiniu. Analizei atlikti naudojamas Andrew Hayes'o sukurtas PROCESS makro modelis (1 modelio konfigūracija), leidžiantis įvertinti moderatoriaus poveikį pagrindinio nepriklausomo kintamojo ir priklausomojo kintamojo ryšiui. Analizės rezultatai pateikiami 36 lentelėje ir 12 priede. Tikrinamos šios hipotezės:

H4a. Dirbtinio intelekto turinio tikrojo autoriaus atskleidimas teigiamai veikia vartotojų, kurie į DI žvelgia kaip į bendraautorį, pasitikėjimą turiniu.

H4b. Dirbtinio intelekto turinio tikrojo autoriaus atskleidimas neigiamai veikia vartotojų, kurie į DI žvelgia kaip į įrankį, pasitikėjimą turiniu.

36 lentelė. Moderacijos analizė tarp autorystės atskleidimo, pasitikėjimo pokyčio ir požiūrio į DI. Modelio tinkamumas (N=461)

Požiūris	Rodiklis	Reikšmė
DI kaip bendraautorius	R	0,4870
	R^2	0,2371
	F reikšmė	47,3550
	p reikšmė	< 0,001
DI kaip įrankis	R	0,4895
	R^2	0,2396
	F reikšmė	47,9956
	p reikšmė	< 0,001

Iš pateikto modelio tinkamumo pagrindimo galima pastebėti, jog abu moderacijos modeliai (su kiekvienu moderatoriumi atskirai) yra statistiškai tinkami – t. y. jų pagalba galima aiškinti pasitikėjimo turiniu pokyčius priklausomai nuo turinio autoriaus ir požiūrio į DI. 37-oje lentelėje siekiama atskleisti ar ši moderacija yra reikšminga.

37 lentelė. Moderacijos analizė tarp autorystės atskleidimo, pasitikėjimo pokyčio ir požiūrio į DI. Rezultatai (N=461)

Požiūris	B	t reikšmė	p reikšmė	Reikšmingumas, kryptis
----------	---	-----------	-----------	------------------------

DI kaip bendraautorius	-0,2321	-2,0895	0,0372	Reikšminga, neigiama
DI kaip įrankis	-0,0511	-0,2988	0,7652	Nereikšminga

Moderacijos analizės rezultatai parodė, kad požiūris į dirbtinį intelektą gali keisti turinio autoriaus atskleidimo poveikį vartotojų pasitikėjimui. Atlikus analizę su moderatoriumi „požiūris į DI kaip į bendraautorius“, nustatyta, kad sąveikos efektas yra statistiškai reikšmingas ($B = -0,2321$, $p = 0,0372$). Tai reiškia, kad kuo labiau respondentai linkę matyti DI kaip kūrybinį partnerį, tuo mažesnis yra turinio autoriaus (žmogus vs DI) įtakos pasitikėjimui skirtumas. Kitaip tariant, teigiamas požiūris į DI kaip į bendraautorius silpnina autoriaus atskleidimo poveikį vartotojų pasitikėjimui.

Tuo tarpu vertinant moderuojančią „požiūrio į DI kaip į įrankį“ įtaką, gautas sąveikos koeficientas buvo neigiamas ($B = -0,0511$), tačiau statistiškai nereikšmingas ($p = 0,7652$). Tai rodo tendenciją, kad kuo labiau respondentai žiūri į DI kaip į paprastą priemonę, tuo labiau neigiamai vertina turinį, kurio autorius – DI. Nors ši tendencija atitinka hipotezės H4b kryptį, statistinių įrodymų jos patvirtinimui nepakanka.

Remiantis moderacijos analizės rezultatais, hipotezę **H4a** (dirbtinio intelekto turinio tikrojo autoriaus atskleidimas neigiamai veikia vartotojų, kurie į DI žvelgia kaip į bendraautorius, pasitikėjimo turiniu pokyčiui) yra **patvirtinama**, o **H4b** (dirbtinio intelekto turinio tikrojo autoriaus atskleidimas teigiamai veikia vartotojų, kurie į DI žvelgia kaip į įrankį, pasitikėjimo turiniu pokyčiui) hipotezę yra **atmetama**.

Siekiant išsiaiškinti, kokią įtaką tikrojo turinio autoriaus identifikavimui daro respondentų žinios apie dirbtinio intelekto įrankius, buvo atliktas chi-kvadrato testas. Šis metodas pasirinktas siekiant identifikuoti ar egzistuoja ryšys tarp pasirinktų kintamųjų ir dėl lengvo rezultatų interpretavimo. Identifikavus statistiškai reikšmingą ryšį tarp kintamųjų, planuota atlikti binarinę logistinę regresinę analizę, tačiau šiuo konkrečiu atveju to neprireikė. Prieš atliekant chi-kvadrato testą, būtina suskirstyti kintamuosius į tinkamas kategorijas. Žinių apie DI kintamasis (Likerto skalė nuo 1 iki 5) buvo suskirstytas į tris kategorijas: mažas žinias (1–2), vidutines (3) ir dideles (4–5). Tuo tarpu turinio autoriaus identifikavimo rezultatai buvo konvertuoti į dvejetainį kintamąjį: 1 = teisingai identifikavo, 0 = neteisingai identifikavo. Atsakymai „3“ (nei sutinku, nei nesutinku) buvo priskirti neteisingai identifikavusiems, kadangi toks pasirinkimas rodo neapsisprendimą ir negebėjimą aiškiai įvertinti turinio autoriaus, todėl negali būti laikomas teisingu identifikavimu. Testo rezultatai detalizuojami 13 priede. Tikrinama ši hipotezė:

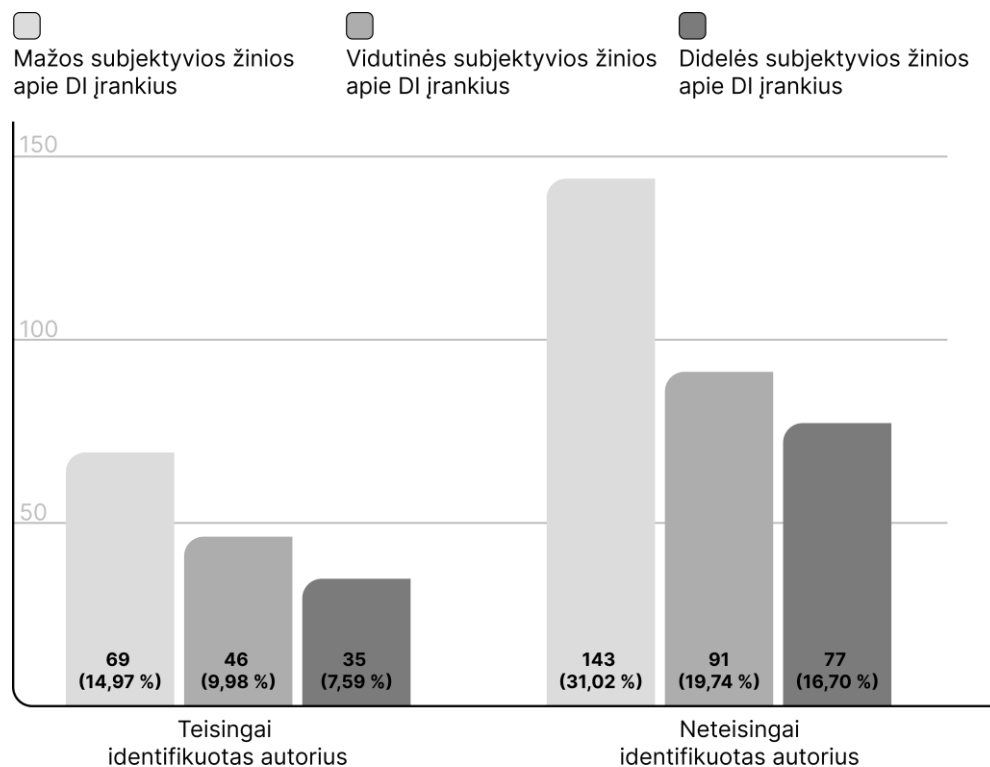
H5. Vartotojų žinios apie DI įrankius teigiamai veikia jų gebėjimą identifikuoti tikrąjį turinio autorių

Toliau pateikiama sistemizuoti chi-kvadrato analizės rezultatai (žr. 38 lentelę). Šiuo atveju priklausomas kintamasis – teisingas autoriaus identifikavimas, o nepriklausomas kintamasis – žinios apie DI įrankius.

38 lentelė. Chi-kvadrato analizė tarp DI žinių kategorijų ir autoriaus identifikavimo ($N=461$)

Pearson Chi-Square (χ^2)	p reikšmė	Cramer's V	Teisingai identifikavo	Neteisingai identifikavo
0,152	0,927	0,018	150	311 (iš jų 183 pažymėjo „nei sutinku, nei nesutinku“ opciją)

Chi-kvadrato analizė parodė, kad nėra statistiškai reikšmingo ryšio tarp respondentų žinių apie DI įrankius ir jų gebėjimo teisingai identifikuoti tikrąjį turinio autorių ($\chi^2 = 0,152$, $p = 0,927$). Efekto dydis pagal Cramer's V buvo labai mažas ($V = 0,018$), o tai rodo, kad ir praktinė sąryšio reikšmė yra nereikšminga. Iš 461 respondentų autorių teisingai identifikavo 32,5 proc. asmenų (150 identifikavo teisingai, o 311 – neteisingai), nepriklausomai nuo jų žinių lygio. Tai reiškia, jog respondentams itin sunku atskirti, kas yra turinio autorius. Toliau pateikiamoje stulpelinėje diagramoje detalizuojami kiekybiniai rezultatai (žr. 6 pav.).



6 pav. Respondentų subjektyvių žinių apie DI įrankius ir faktinio autoriaus identifikavimo kryžminė analizė (N=461)

Respondentų gebėjimas teisingai identifikuoti turinio autorių nedemonstravo aiškos priklausomybės nuo jų deklaruotų žinių apie dirbtinio intelekto įrankius. Tiek mažas, tiek vidutines ar dideles žinias deklaravę asmenys panašiai pasiskirstė tarp tų, kurie teisingai ir neteisingai identifikavo autorių. Didžioji dalis respondentų – nepriklausomai nuo žinių lygio – neteisingai identifikavo autorių, o tai išryškina bendrą tendenciją: aukštesnis žinių lygis nereiškia tikslesnio autoriaus identifikavimo. Tai pat pastebima, kad net 46 proc. respondentų pripažino, jog apie DI įrankius žino nedaug. Svarbu atkreipti dėmesį ir į tai, jog asmenys savo žinias apie DI įrankius vertino subjektyviai, o tai gali turėti įtakos galutiniam rezultatui.

Remiantis chi-kvadrato testo rezultatais, hipotezė **H5** (vartotojų žinios apie DI įrankius teigiamai veikia jų gebėjimą identifikuoti tikrąjį turinio autorių) yra **atmetama**.

Siekiant nustatyti kokią įtaką respondentų subjektyvus gebėjimas lengvai atskirti DI turinį nuo žmogaus sukurto daro jo faktiniam gebėjimui tai atlikti, įvykdytas chi-kvadrato testas, vadovaujantis panašia logika, aptarta analizuojant H5. Chi-kvadrato testui identifikavus statistiškai reikšmingą ryšį, buvo planuojama atlikti ir binarinę logistinę regresinę analizę ryšio kryptčiai ir įtakai nustatyti, tačiau ir šįkart jos neprireikė. Prieš atliekant chi-kvadrato testą, Likerto skale matuojamas kintamasis, skirtas

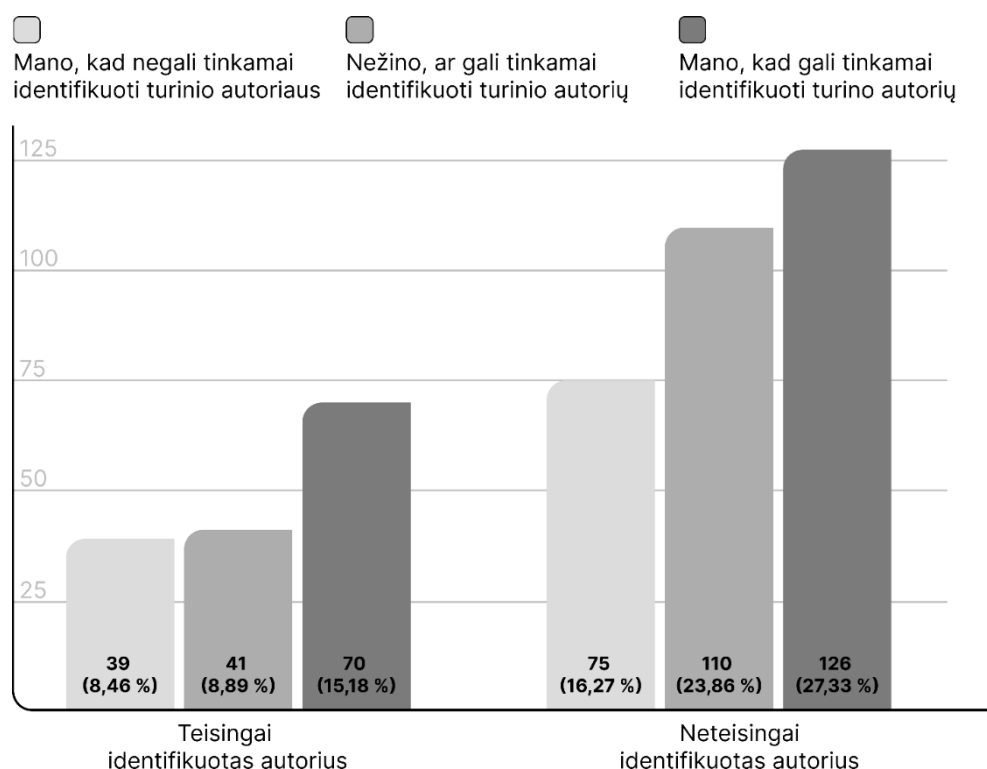
įvertinti respondentų subjektyvų gebėjimą atskirti DI turinį nuo žmogaus, buvo suskaidytas į tris kategorijas (1 – mano, kad negali tinkamai identifikuoti turinio autoriaus; 2 – nežino ar gali tinkamai identifikuoti turinio autorių; 3 – mano, kad gali tinkamai identifikuoti turinio autorių). Testo rezultatai detalizuojami 13 priede ir 39 lentelėje. Tikrinama ši hipotezė:

H6. Vartotojų suvokiamas gebėjimas tinkamai identifikuoti DI turinį teigiamai veikia jų gebėjimą identifikuoti DI turinį.

39 lentelė. Chi-kvadrato analizė tarp subjektyvaus gebėjimo identifikuoti turinio autorių ir autoriaus identifikavimo (N=461)

Pearson Chi-Square (χ^2)	p reikšmė	Cramer's V
3,041	0,219	0,081

Chi-kvadrato analizė parodė, kad suvokiamas gebėjimas atpažinti DI turinį statistiškai reikšmingo ryšio su tikru autoriaus identifikavimo tikslumu neturi ($\chi^2 = 3,041$, $p = 0,219$). Efekto dydis, vertinamas Cramer's V (0,081), rodo labai silpną sąsają tarp kintamųjų. Toliau pateikiami kiekybiniai duomenys (žr. 7 pav.).



7 pav. Respondentų suvokiamo gebėjimo atskirti DI turinį nuo žmogaus kurto turinio ir faktinio autoriaus identifikavimo kryžminė analizė (N=461)

Iš stulpelinės diagramos matyti, kad daugiausia respondentų tiek tarp teisingai, tiek tarp neteisingai identifikavusiųjų turinio autorių nurodė, jog mano galintys tinkamai atpažinti turinio autorių – iš visos imties net 42,5 proc. respondentų manė, jog tinkamai identifikuoti turinio autorių galėtų. Tačiau net ir tarp jų neteisingai identifikavusiųjų yra beveik dvigubai daugiau (126) nei teisingai identifikavusiųjų (70). Tai rodo, kad subjektyvus pasitikėjimas savo gebėjimu identifikuoti tikrąjį turinio autorių neatitinka realaus gebėjimo tai padaryti.

Remiantis chi-kvadrato testo rezultatais, hipotezė **H6** (vartotojų suvokiamas gebėjimas tinkamai identifikuoti DI turinį teigiamai veikia jų gebėjimą identifikuoti DI turinį) yra **atmetama**.

Siekiant įvertinti, ar vartotojų subjektyvus suvokimas, jog žinojimas apie tikrąjį turinio autorių didina jų pasitikėjimą turiniu, daro teigiamą įtaką faktiniam pasitikėjimo padidėjimui sužinojus tikrąjį turinio autorių, buvo atlikta tiesinės regresijos analizė. Analizės rezultatai pateikiami 40-oje lentelėje ir 11 priede. Tikrinama ši hipotezė:

H7. Vartotojų suvokiamas pasitikėjimo turiniu didėjimas žinant autorių teigiamai veikia jų pasitikėjimą turiniu, atskleidus tikrąjį turinio autorių.

40 lentelė. Paprastoji regresija tarp suvokiamo pasitikėjimo didėjimo sužinojus autorių ir faktinio pasitikėjimo pokyčio. Modelio tinkamumas (N=461)

Rodiklis	Reikšmė
R	0,077
R ²	0,006
F reikšmė	2,754
p reikšmė	0,098

Tiesinės regresijos analizė parodė, kad tarp vartotojų subjektyvaus pasitikėjimo turiniu, kai žinomas jo autorius, ir faktinio pasitikėjimo turiniu padidėjimo sužinojus autorių, nėra statistiškai reikšmingo ryšio ($F = 2,754$, $p = 0,098$). Tai patvirtina ir $R^2 = 0,006$ rodiklis, kurio reikšmė suponuoja, kad modelis paaiškina 0,6 % duomenų. Nors regresijos koeficientas $B = -0,081$ (žr. 11 priedą) rodo nežymią neigiamą tendenciją, šis efektas nėra reikšmingas, todėl negalima daryti tvirtų išvadų apie šio ryšio egzistavimą.

Remiantis tiesinės regresijos analizės rezultatais, hipotezė **H7** (vartotojų suvokiamas pasitikėjimo turiniu didėjimas žinant autorių teigiamai veikia jų pasitikėjimą turiniu, atskleidus tikrąjį turinio autorių) yra **atmetama**.

Siekiant įvertinti, ar turinio autoriaus (žmogus ar dirbtinis intelektas) daro reikšmingą poveikį vartotojų suvokiamam turinio autentiškumui ir kokybei, pasitelkiamas Mann–Whitney U testas, leidžiantis palyginti dviejų nepriklausomų grupių duomenis, kai jie neatitinka normalaus skirstinio. Testo rezultatai pateikiami 41-oje lentelėje ir 10 priede. Tikrinamos šios hipotezės:

H8a. Žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnį teigiamą poveikį vartotojų suvokiamam turinio veržlumui ir įtaigumui nei dirbtinio intelekto sukurtas turinys

H8b. Žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnį neigiamą poveikį vartotojų suvokiamam turinio keistumui nei dirbtinio intelekto sukurtas turinys

H8c. Žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnį neigiamą poveikį vartotojų suvokiamam turinio sintetiškumui nei dirbtinio intelekto sukurtas turinys

41 lentelė. Neparametrinis Mann-Whitney U testas tarp turinio autoriaus atskleidimo bei autentiškumo ir kokybės dimensijų (N=461)

Autentiškumo ir kokybės dimensija	Vidutinis rangas (M)		Z reikšmė	p reikšmė
	Autorius – DI (N=232)	Autorius – žmogus (N=229)		

Veržlumas ir įtaigumas	258,81	202,83	-4,516	<0,001
Keistumas	219,80	242,35	-1,833	0,067
Sintetiškumas	213,62	248,60	-2,846	0,004

Veržlumo ir įtaigumo sąsajos su turinio autoriumi rezultatai rodo statistiškai reikšmingą skirtumą tarp grupių (M (DI) = 258,81, M (žmogus) = 202,83, Z = -4,516, p < 0,001). Vidutinis rangas buvo mažesnis žmogaus turinio scenarijuje (M = 202,83), palyginti su DI scenarijumi (M = 258,81). Tai leidžia teigti, kad dirbtinio intelekto kurtas turinys vartotojams atrodė kokybiškesnis ir autentiškesnis nei žmogaus. Keistumo komponento analizė parodė statistiškai nereikšmingą skirtumą tarp grupių (p = 0,067), nors ir buvo pastebėta tendencija, kad žmogaus turinys buvo vertinamas kaip šiek tiek keistesnis. Sintetiškumo komponentas parodė reikšmingą skirtumą tarp grupių (M (DI) = 213,62, M (žmogus) = 248,60; Z = -2,846; p = 0,004). Didesnis rangas žmogaus turinio grupėje (M = 248,60) leidžia manyti, kad žmogaus kurtas turinys buvo vertinamas kaip labiau sintetinis nei DI sukurtas. Visi šie rezultatai kartu leidžia susidaryti nuomonę, kad respondentai ne tik netinkamai identifikavo turinio autorių (kaip buvo atskleista aukščiau pateiktoje kryžminėje analizėje), bet ir dirbtinio intelekto sukurtą turinį laikė autentiškesniu ir kokybiškesniu nei žmogaus sukurtą.

Remiantis Mann-Whitney U testo rezultatais hipotezės **H8a** (žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio veržlumui ir įtaigumui nei dirbtinio intelekto sukurtas turinys), **H8b** (žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnę neigiamą poveikį vartotojų suvokiamam turinio keistumui nei dirbtinio intelekto sukurtas turinys), **H8c** (žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnę neigiamą poveikį vartotojų suvokiamam turinio sintetiškumui nei dirbtinio intelekto sukurtas turinys) yra **atmetamos**.

Siekiant įvertinti, ar turinio autoriaus (žmogus ar dirbtinis intelektas) daro reikšmingą poveikį vartotojų pasitikėjimui turiniu, naudojamas Mann-Whitney U testas, leidžiantis palyginti dviejų nepriklausomų grupių duomenis, kai jie neatitinka normalaus skirstinio. Testo rezultatai pateikiami 42-oje lentelėje ir 10 priede. Tikrinama ši hipotezė:

H9. Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų pasitikėjimui turiniu nei dirbtinio intelekto sukurtas turinys.

42 lentelė. Neparametrinis Mann-Whitney U testas tarp turinio autoriaus bei pasitikėjimo turiniu ($N=461$)

Rodiklis	Reikšmė
N (autorius – DI / žmogus)	232 / 229
Vidutinis rangas (autorius – DI)	252,37
Vidutinis rangas (autorius – žmogus)	209,35
Z reikšmė	-3,653
p reikšmė	< 0,001

Turinio autoriaus sąsajos su pasitikėjimu turiniu rodo statistiškai reikšmingą skirtumą tarp grupių (M (DI) = 252,37, M (žmogus) = 209,35, Z = -3,653, p < 0,001). Vidutinis rangas buvo mažesnis žmogaus turinio scenarijuje (M = 209,35), palyginti su DI scenarijumi (M = 252,37). Tai reiškia, jog, priešingai nei buvo prognozuota tyrimo literatūrinės analizės dalyje, respondantai dirbtinio intelekto kurtą turinį suvokė kaip patikimesnį, nei žmogaus sukurtą, kai jiems neatskleista, kas yra tikrasis autorius.

Įvertinus Mann-Whitney U testo rezultatus, hipotezė **H9** (žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų pasitikėjimui turiniu nei dirbtinio intelekto sukurtas turinys) yra **atmetama**.

Apibendrinant duomenų analizės ir hipotezių tikrinimo rezultatus, pateikiami tyrimo hipotezių rezultatai (žr. 43 lentelę).

43 lentelė. Tyrimo hipotezių rezultatai

Nr.	Hipotezė	Rezultatas
H1	Vartotojų suvokiamas turinio autentiškumas ir kokybė teigiamai veikia jų pasitikėjimą turiniu.	
	H1a Vartotojų suvokiamas turinio veržlumas ir įtaigumas daro teigiamą įtaką turinio autentiškumo ir kokybės vertinimui.	Patvirtina
	H1b Vartotojų suvokiamas turinio keistumas daro neigiamą įtaką turinio autentiškumo ir kokybės vertinimui.	Atmesta
	H1c Vartotojų suvokiamas turinio sintetiškumas daro neigiamą įtaką turinio autentiškumo ir kokybės vertinimui.	Patvirtinta
H2	Tikrojo turinio autoriaus atskleidimas moderuoja pasitikėjimą turiniu.	
	H2a Dirbtinio intelekto kurto turinio autoriaus atskleidimas daro neigiamą įtaką pasitikėjimui turiniu.	Patvirtinta
	H2b Žmogaus kurto turinio autoriaus atskleidimas daro teigiamą įtaką pasitikėjimui turiniu.	Patvirtinta
H3	Vartotojų pasitikėjimo turiniu pokytis teigiamai veikia jų polinkį sąveikauti su turiniu.	Sąlyginai patvirtina*
H4	Požiūris į DI moderuoja kaip vartotojų pasitikėjimo pokytį veikia tikrojo turinio autoriaus atskleidimas.	
	H4a Dirbtinio intelekto turinio tikrojo autoriaus atskleidimas neigiamai veikia vartotojų, kurie į DI žvelgia kaip į bendraautorių, pasitikėjimo turiniu pokyčiui.	Patvirtinta
	H4b Dirbtinio intelekto turinio tikrojo autoriaus atskleidimas teigiamai veikia vartotojų, kurie į DI žvelgia kaip į įrankį, pasitikėjimo turiniu pokyčiui.	Atmesta
H5	Vartotojų žinios apie DI įrankius teigiamai veikia jų gebėjimą identifikuoti tikrąjį turinio autorių.	Atmesta
H6	Vartotojų suvokiamas gebėjimas tinkamai identifikuoti DI turinį teigiamai veikia jų gebėjimą identifikuoti DI turinį.	Atmesta
H7	Vartotojų suvokiamas pasitikėjimo turiniu didėjimas žinant autorių teigiamai veikia jų pasitikėjimą turiniu, atskleidus tikrąjį turinio autorių.	Atmesta
H8	Žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio autentiškumui ir kokybei nei dirbtinio intelekto sukurtas turinys.	
	H8a Žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio veržlumui ir įtaigumui nei dirbtinio intelekto sukurtas turinys.	Atmesta
	H8b Žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnę neigiamą poveikį vartotojų suvokiamam turinio keistumui nei dirbtinio intelekto sukurtas turinys.	Atmesta
	H8c Žmogaus sukurtas turinys socialiniuose tinkluose daro stipresnę neigiamą poveikį vartotojų suvokiamam turinio sintetiškumui nei dirbtinio intelekto sukurtas turinys.	Atmesta
H9	Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų pasitikėjimui turiniu nei dirbtinio intelekto sukurtas turinys.	Atmesta

*paaiškinimas po 34 lentele.

Šiame baigiamojo darbo skyriuje buvo išsamiai nagrinėti pagrindinių tyrimo kintamųjų ryšiai ir tikrintos tyrimo hipotezės, siekiant geriau suprasti vartotojų pasitikėjimo dirbtinio intelekto sukurtu turiniu dinamiką. Naudojant koreliacinę analizę, identifikuoti reikšmingi ryšiai tarp tokių kintamųjų kaip įtaigumas, keistumas, pasitikėjimas, suvoktas DI turinio sintetiškumas ar autoriaus atpažįstamumas. Duomenų analizė atskleidė, kad emociniai turinio aspektai – įtaigumas ir veržlumas – stipriausiai veikia pasitikėjimą turiniu bei respondentų požiūrį į DI kaip į bendraautorių. Tuo tarpu didesnis sintetiškumo pojūtis dažniausiai susijęs su sumažėjusiu pasitikėjimu ir mažesniu noru sąveikauti su turiniu. Nors kai kurios sąsajos tarp kintamųjų buvo silpnos, jų statistinis reikšmingumas rodo, kad net sudėtingose temose nedideli ryšiai gali būti svarbūs.

Statistinių testų (regresinės analizės, ANOVA, Mann–Whitney U testų, chi-kvadrato analizės, moderacijos analizės) pagalba patvirtintos 6 iš 14 hipotezių. Patvirtinimai rodo, kad įtaigumas ir veržlumas teigiamai veikia pasitikėjimą turiniu (H1a), sintetiškumas – neigiamai (H1c), tikrojo autoriaus atskleidimas turi reikšmingą įtaką pasitikėjimui (H2a, H2b), pasitikėjimas turiniu skatina sąveiką su juo (H3), teigiamas požiūris į DI kaip bendraautorių švelnina autoriaus atskleidimo poveikį pasitikėjimui turiniu (H4a).

Tuo tarpu dauguma kitų hipotezių, susijusių su vartotojų žiniomis, suvokimu ar žmogaus sukurtu turinio poveikiu, buvo atmestos. Pavyzdžiui, nei vartotojų žinios apie DI, nei jų subjektyvi nuomonė, kad geba atpažinti DI turinį, nedarė reikšmingos įtakos faktiniam autoriaus identifikavimui (H5, H6). Taip pat paaiškėjo, kad žmonės dažnai klaidingai identifikuoja turinio autorių, o DI kurtas turinys neretai vertinamas kaip autentiškesnis ir įtaigesnis nei žmogaus kurtas (atmestos H8 hipotezės).

4.5. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose įtakos vartotojų pasitikėjimui bei ketinimams tyrimo rezultatų apibendrinimas, mokslinė diskusija, tyrimo ribotumai bei tolimesnių tyrimų kryptys

Siekiant palyginti dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikį vartotojų pasitikėjimui bei ketinimams lyginamosios analizės metodu, buvo sudarytas conceptualusis modelis, paremtas mokslinės literatūros analize. Vadovaujantis modeliu bei mokslininkų išvalgomis, buvo sudaryti 6 tyrimo uždaviniai bei tyrimo modelis. Pagal pastarąjį buvo iškelta 11 hipotezių – 6 iš jų buvo detalizuotos į keletą sub-hipotezių. Pastarosios suteikė pagrindą tyrimo įrankiui – internetinei apklausai – sudaryti. Visgi, po faktorinės analizės buvo nuspręstą tyrimą tęsti su 9 hipotezėmis. Toliau pateikiamos bendros tyrimo išvalgos, jų suderinamumas su teorine literatūra, aptariami tyrimo ribotumai bei tolimesnių tyrimų kryptys.

Tyrimo pagalba buvo siekiama išsiaiškinti, kokią įtaką vartotojų žinios apie dirbtinio intelekto įrankius daro jų gebėjimui tinkamai identifikuoti turinio autorių. Remiantis gautais rezultatais, neaptikta statistiškai reikšmingo ryšio tarp respondento deklaruojamų žinių apie DI įrankius ir jo faktinio gebėjimo teisingai identifikuoti tikrąjį turinio autorių. Respondentai, kurie nurodė turintys daug žinių apie dirbtinio intelekto įrankius, ne ką dažniau atspėjo, kas sukūrė pateiktą įrašą, palyginus su tais, kurie teigė mažiau nusimanantys šioje srityje. Socialinių tinklų naudotojai Lietuvoje, net ir turintys daugiau žinių apie DI ne visada geba tinkamai identifikuoti turinio autorių. Taip pat rezultatai atskleidė, kad vos 32,5 proc. tyrimo dalyvių tinkamai identifikavo turinio autorių, tuo tarpu net 39,7 proc. dalyvių neturėjo tvirtos nuomonės apie turinio autorystę. Iš to galima daryti prielaidą, kad DI turinys tapo toks įtikinamas, kad geba suklaidinti net ir patyrusius vartotojus. Park et al. (2024) teigia, jog dabartinėje skaitmeninėje erdvėje vis sunkiau atpažinti turinio autorystę ir užtikrinti skaidrumą –

DI generuojami tekstai ir vaizdai kokybės aspektu kryptingai artėja prie žmogiškosios kūrybos kokybės. Larsson'as ir Heintz'as (2020) nurodo, kad DI turiniui dažnai trūksta autentiškumo požymių, tačiau eiliniam vartotojui gali būti sudėtinga tuos požymius identifikuoti. Šio eksperimentinio tyrimo dalyviai turinio autorystę turėjo nustatyti tik iš pateiktos tekstinės ir vizualinės informacijos, be kitų užuominų apie turinio autorių. Kaip pažymi Jenkins et al. (2020), vartotojai, kai autorius nėra aiškiai nurodytas, remiasi periferiniais signalais – teksto tonu, stiliumi, vizualine kokybe – tačiau šie spėjimai gali būti klaidingi. Papildomai buvo vertintas ir subjektyvus gebėjimas atskirti DI turinį – t. y. ar respondentai patys mano sugebantys atpažinti DI kurtus įrašus. Ir šiuo atveju rezultatai parodė, kad pasitikėjimas savo jėgomis nieko nereiškia – manantys gebantys atpažinti DI turinį neidentifikavo jo tiksliau nei tie, kurie abejojo savo gebėjimais. Lietuvos gyventojams, nepriklausomai nuo jų deklaruojamų žinių apie dirbtinį intelektą ar gebėjimo atpažinti DI sukurtą turinį, dažnai kyla sunkumų teisingai identifikuojant tokį turinį. Tai rodo, kad šiuo klausimu efektyviausiai galėtų padėti kiti DI pagrindu sukurti įrankiai, kurie specializuojasi aptikdami šį turinį. Tyrimas taip pat atskleidė, kad jaunesni žmonės ir tie, kurie dažniau naudojami socialiniais tinklais, yra linkę geriau vertinti savo gebėjimą atskirti DI sukurtą turinį nuo žmogaus sukurto. Panaši tendencija pastebima ir kalbant apie jų žinias apie DI įrankius – jaunesni, daugiau laiko socialiniuose tinkluose praleidžiantys ir dažniau DI įrankius naudojančios žmonės nurodo, kad apie šias technologijas žino daugiau nei kiti.

Tyrimo rezultatai taip pat parodė, kad turinio autoriaus atskleidimas daro reikšmingą įtaką vartotojų pasitikėjimui: atskleidus, jog turinį sukūrė dirbtinis intelektas, pasitikėjimas juo smuko, kitą vertus, jei tikrasis turinio autorius buvo žmogus, pasitikėjimas juo didėjo. Šis dėsningumas patvirtina, jog auditorija socialiniuose tinkluose labiau linkusi pasitikėti turiniu, kai žino, kas jo autorius – žmogus, o ne DI sprendimas. Panašią tendenciją išvėlgė ir Buchanan'as bei Hickman'as (2024), nustatydami, kad vien informacija apie DI indėlį turinio kūrimo gali sumažinti skaitytojų pasitikėjimą žinute. Jų eksperimente naujienų straipsniuose paminėjus, kad turinį generavo DI, skaitytojai iškart vertino jį skeptiškiau, net jei tiksliai negalėjo įvardinti DI indėlio apimtį. Tai atspindi visuomenėje paplitusią nuostatą, jog žmogaus parašytas turinys yra patikimesnis už DI sugeneruotą, nes, vartotojų suvokimu, žmogus laikomas atsakingu už savo darbą, o DI yra naujovė, galinti daryti klaidas. Net esant objektyviam ir faktiniam DI turiniui, jam dažnai trūksta emocinio ryšio, kuris padidina pasitikėjimą – vartotojai labiau įsitraukia į žmogaus kurtą įrašą dėl jo žmogiško tono (Belanche et al., 2024). Visgi literatūroje pastebimas ir niuansas – kai kurių autorių teigimu, vien turinio paženklinimas kaip „DI generuoto“ nebūtinai visada mažina pasitikėjimą. Pavyzdžiui, Li ir Yang'o (2024) eksperimento metu aiškus DI turinio žymėjimas neturėjo reikšmingos įtakos vartotojų suvoktam žinutės patikimumui ar ketinimui ja dalintis. Šis skirtumas nuo Buchanan'o ir Hickman'o (2024) išvadų leidžia manyti, jog turinio kontekstas bei pateikimo būdas gali lemti, ar DI autorystės atskleidimas sumažins pasitikėjimą. Šio tyrimo atveju, autorius atskleistas tik po to, kai respondentai jau buvo įvertinę įrašo patikimumą. Atskleidus autorių, respondentų pasitikėjimas turiniu buvo vertinimas dar kartą. Šis metodas leido tiesiogiai palyginti kokią įtaką turinio autoriaus atskleidimas daro pasitikėjimui turiniu: nustatyta, kad sužinojus, jog turinio autorius yra žmogus, Lietuvoje gyvenantys socialinių tinklų naudotojai jį įvertino patikimiau nei tada, kai paaiškėjo, jog tikrasis turinio autorius buvo dirbtinis intelektas. Ši tendencija pastebima nepriklausomai nuo respondentų amžiaus, lyties, išsilavinimo, socialinių tinklų naudojimo dažnumo bei trukmės, dirbtinio intelekto įrankių naudojimo dažnio.

Tyrimo duomenys atskleidė, kad vartotojų suvokiamas turinio autentiškumas ir kokybė yra glaudžiai susiję su jų pasitikėjimu turiniu. Atlikus regresinę analizę nustatyta, jog turinys, kuris respondentams

atrodo įtaigus ir veržlus, labiau skatina jų pasitikėjimą. Kitaip tariant, jeigu socialinio tinklo įrašas vertinamas kaip sąžiningas, nuoseklus, emocionaliai paveikus, įtraukiantis, aktualus ir suprantamas, vartotojai yra linkę juo labiau pasikliauti. Visgi, verta pastebėti, kad, atlikta faktorinė analizė parodė, jog respondentai pastarųjų autentiškumo ir kokybės dimensijų neišskyrė individualiai, o yra linkę vertinti jas bendrai kaip teigiamas turinio savybes. Mokslinėje literatūroje išskiriama, kad vartotojai kur kas labiau pasitiki turiniu, kuris jiems atrodo autentiškas, t. y. nuoseklus bei emociškai įtraukiantis, ir kokybiškas (Eigenraam et al., 2021; Kim et al., 2015; J. Yang et al., 2021). Priešingai, jeigu turinys atrodo nenatūralus, prieštaringas arba nekokybiškas, vartotojų pasitikėjimas juo krenta. Ismagilova et al. (2020) bei Mirbabaie ir Stieglitz'as (2021) pažymi, kad net menkiausi neatitikimai turinyje ar jo pateikime, pvz., matomi šeši rankos pirštai, kai tikimasi penkių, gali signalizuoti neautentiškumą ir galimą manipuliaciją, dėl ko vartotojai praranda pasitikėjimą. Šio tyrimo kontekste taip pat išryškėjo tendencija, kad tokie aspektai kaip turinio „sintetiškumas“ gali neigiamai veikti bendrą patikimumo vertinimą. Tuo tarpu respondentų suvokiamas turinio „keistumas“ statistiškai reikšmingos įtakos nedaro – šiuo atveju gali būti, jog kai kurie respondentai jį gali sieti su turinio novatoriškumu ir inovatyvumu. Visgi, pastebėta tendencija, kad dirbtinį intelektą dažniau naudojantys respondentai yra linkę turinį vertinti kaip labiau įtaigesnį ir veržlesnį ir šiek tiek mažiau keistą. Pažymėtina, kad bendrai imties dalyvių vertinimai turinio autentiškumo ir kokybės klausimu buvo ne vienareikšmiški: bendrai turinio „įtaigumo ir veržlumo“ vertinimo vidurkiai sukosi aplink vidurį, t. y. opcija „nei sutinku, nei nesutinku“, panaši tendencija įžvelgiama ir „sintetiškumo“ vertinime, visgi, turinio „keistumo“ vidurkis buvo 2,6, o tai rodo jog respondentai turinį buvo linkę vertinti kaip mažiau keistą. Tai rodo tam tikrą skeptiškumą: vartotojai nėra linkę kategoriškai įvertinti turinio priskirti aukšto autentiškumo nepažįstamam turiniui. Vis dėlto net ir esant vidutiniams vertinimams, ryšys išryškėjo: tie, kurie jautėsi labiau įtikinti turinio arba pastebėjo jo aukštesnę kokybę, demonstravo didesnę pasitikėjimą.

Šio konkretaus eksperimentinio tyrimo atveju, respondentai žmogaus kurtą turinį įvertino mažiau kokybišką ir autentišką nei dirbtinio intelekto. Nors literatūroje paprastai pabrėžiama, kad vartotojai labiau pasitiki žmogaus kurtu turiniu dėl didesnio jo įtaigumo ir veržlumo bei mažesnio „sintetiškumo“ (Huschens et al., 2023; Buchanan ir Hickman, 2024; Eigenraam et al., 2021), šio tyrimo rezultatai atskleidė priešingą tendenciją. Respondentų vertinimu, dirbtinio intelekto sugeneruotas turinys pasirodė esąs labiau įtaigus ir veržlus nei žmogaus sukurtas, be to, jam priskirtas mažesnis „sintetiškumo“ jausmas. Toks rezultatas atrodo paradoksalus turint omenyje, jog autentiškumas laikomas esmine pasitikėjimo sąlyga (Eigenraam et al., 2021) ir paprastai siejamas su žmogaus kūryba. Vienas galimas šio neatitikimo paaiškinimas – sparčiai didėjanti DI generuojamo turinio kokybė. Pavyzdžiui, Park et al. (2024) nustatė, kad vartotojams darosi sunku atskirti aukštos kokybės DI sukurtą turinį nuo žmogaus turinio, o pastarojo įtikinamumas ir kokybė gali būti vertinami panašiai. Kai DI turinys yra pakankamai sklandus, emociškai įtraukiantis ir neturi akivaizdžių „sintetiškumo“ ženklų, vartotojai gali jį suvokti kaip autentišką ir patikimą. Šio tyrimo atvejis rodo, kad turinio sąžiningumas, nuoseklumas, aktualumas, suprantamumas, gebėjimas užmegzti emocinį ryšį ir įtraukti gali nusverti išankstines nuostatas – jei DI sukurtas pranešimas sugeba sudaryti autentiškumo įspūdį ir atrodyti kokybišku, vartotojų pasitikėjimas juo gali prilygti, ar net pranokti pasitikėjimą žmogaus kurtu turiniu.

Analizuojant vartotojų elgseną, tyrimo rezultatai patvirtino, kad pasitikėjimas turiniu daro teigiamą įtaką vartotojų ketinimams sąveikauti su tuo turiniu. Didesnis pasitikėjimo lygis nežymiai, bet statistiškai reikšmingai padidina vartotojo ketinimą sąveikauti su turiniu: paspausti „patinka“,

komentuoti, pasidalinti ar apsilankyti autoriaus paskyroje. Ši tendencija atitinka teorinius modelius: moksliniai tyrimai patvirtina, kad pasitikėjimas yra būtina sąlyga vartotojų ketinimui sąveikauti su turiniu. Autoriai išskiria, kad pasitikėjimas informacijos šaltiniu didina informacijos priėmimą – tuomet vartotojas linkęs vykdyti šaltinio rekomendacijas, dalintis turiniu ar net įsigyti rekomenduojamą prekę (Almira Vania & Daniel, 2023; Dabbous & Barakat, 2020). Priešingai, jei vartotojas nepasitiki įrašu, jis greičiausiai ignoruos turinio siūlymą veikti – nekomentuos, nesidalins ir nereaguos (N. N. Mohamed et al., 2023b). Tyrimo duomenys šią tendenciją patvirtina: pasitikėjimas yra veikia vartotojo ketinimą veikti. Vis dėlto svarbu paminėti ir tai, kad, nepaistant ryšio tarp pasitikėjimo ir ketinimų statistiško reikšmingumo, jis yra silpnas. Pasitikėjimas paaiškina tik apie 4 % vartotojų ketinimo sąveikauti su turiniu. Tai rodo, kad be pasitikėjimo egzistuoja ir kitų veiksnių, lemiančių vartotojų sprendimą sureaguoti į turinį. Dėl šios priežasties pasitikėjimas traktuojamas kaip būtina, bet nepakankama sąlyga: net ir pasitikėjimą pelnęs turinys nebūtinai garantuos įsitraukimą, jeigu, pavyzdžiui, pats turinys neskaitins veiksmo, ar bus neaktualus. Shahbaznezhad'as (2021) išskiria, kad priklausomai nuo turinio tipo, galima tikėtis skirtingo įsitraukimo: transakcinio ir emocinio tipo įrašai skatina didesnę įsitraukimą, tuo tarpu racionaliujų įrašų pagrindinis tikslas informuoti ir pastarieji paprastai sulaukia mažesnio įsitraukimo. Šio tyrimo atveju buvo pasirinktas racionalaus tipo įrašas, o tai galėjo lemti sąlyginai mažą ketinimą įsitraukti, nepriklausomai nuo pasitikėjimo turiniu lygio. Visgi, pastebėta, jog pasitikėjimas turiniu skatina Lietuvoje gyvenančių lietuvių įsitraukimą, tačiau, atsižvelgiant į ribotą ryšio stiprumą, tikėtina, jog vartotojų ketinimus lemia ir kiti elementai – emocinis įsitraukimas, skatinimas veikti –, kurie gali būti pagrindu tolimesniems tyrimams šia tema. Taip pat pastebėta, jog respondentų lytis ir amžius daro reikšmingą įtaką ketinimui sąveikauti su turiniu: moterys ir vyresnio amžiaus asmenys yra labiau linkę įsitraukti į turinį.

Galiausiai tyrime nagrinėtas vartotojų požiūris į dirbtinį intelektą kaip galimas moderatorius, darantis tiesioginę įtaką turinio autorystės atskleidimo poveikiui pasitikėjimu turiniu. Remiantis teorinėmis įžvalgomis, buvo tikimasi, kad respondentų, kurie į DI žvelgia palankiau – šio tyrimo kontekste kaip į bendraautorį – pasitikėjimas turiniu pakis mažiau nei tų, kurių požiūris į DI yra neigiamas. Moderacijos rezultatai išryškino, kad ši prielaida iš dalies yra teisinga: palankiau į DI žvelgiančių asmenų pasitikėjimo pokytis buvo reikšmingai mažesnis atskleidus tikrąjį turinio autorių, nors šis ryšys buvo gan silpnas. Visgi, nebuvo atskleistas reikšmingas pasitikėjimo pokytis tuo atveju, jei žmogus į DI žvelgia neigiamiau. Svarbu paminėti, kad net ir teigiamai DI atžvilgiu nusiteikusių respondentų pasitikėjimas pastebimai sumažėdavo sužinojus, kad turinį sukūrė DI. Taip pat, nors kai kurie respondentai subjektyviai manė, kad jų pasitikėjimas turiniu padidėtų sužinojus tikrąjį jo autorių, praktikoje ši tendencija nebuvo patvirtinta. Buchanan'as ir Hickman'as (2024) savo tyrime užsiminė apie išankstinį visuomenės nusistatymą: nepaisant to, kad DI sprendimai populiarėja, egzistuoja įsišaknijęs įsitikinimas, kad žmogus yra patikimesnis informacijos kūrėjas. Tyrimo dalyviai, kurie į DI žvelgia kaip į naudingą įrankį ar bendraautorį, teoriškai turėjo būti atviresni DI turiniui (Sargin, 2024), tačiau praktikoje jų vertinimai vis tiek pablogėdavo, kai paaiškėdavo, kad tekstą parašė ne žmogus. Taip galėjo nutikti dėl kognityvinio disonanso tarp racionalaus ir emocinio vertinimo: vartotojas gali racionaliai pripažinti DI privalumus (jo greitį, efektyvumą), tačiau emociniame lygmenyje vis tiek stokoti pasitikėjimo DI kūryba. Scharowski'o et al. (2023) tyrimas iš dalies tai patvirtina – net ir pateikus vartotojams išsamius, į žmogų orientuotus DI sprendimų paaiškinimus, taip didinant skaidrumą, pasitikėjimas DI sistema ne visada padidėja. Vien skaidrumas ar palankus požiūris nebūtinai lemia didesnę pasitikėjimą, kai reikia vertinti konkretų DI darbo rezultatą. Tam įtakos gali turėti ir pasąmonėje glūdintis supratimas, kad žmogus turi intenciją, moralę

ir atsakomybę, o DI – paremtas algoritmais ir neatsako už klaidingai pateiktą informaciją. Taip pat pastebėta, jog vyresnio amžiaus asmenys yra šiek tiek labiau linkę į DI žvelgti kaip į bendraautorį, nei jaunesni. Bendrai rezultatai leidžia daryti prielaidą, kad palankus požiūris į DI sumažina, o ne visiškai eliminuoja neigiamą DI autorystės poveikį pasitikėjimui.

Atliekant tyrimą susidurta su keliais **tyrimo ribotumais**, galinčiais turėti įtakos rezultatų interpretacijai:

1. Tyrimo imtis buvo atrinkta patogumo atrankos būdu – nors dalyvių skaičius buvo pakankamas (N=461) ir demografiškai įvairus, respondentai nusprendė savanoriškai prisijungti prie apklausos. Tai riboja rezultatų apibendrinimą platesniam kontekstui: kitose šalyse ar kultūrose vartotojų reakcijos į DI ir žmogaus turinį gali skirtis.
2. Tyrime naudotas scenarijais grįstas eksperimentas, pasitelkiant internetinę apklausą. Šiuo atveju kyla validumo klausimas – realioje socialinių tinklų aplinkoje vartotojų elgesys gali būti kitoks nei eksperimento metu, kur dalyvis žino, jog vertina pateiktą turinį tyrimo tikslais. Pavyzdžiui, realiame „Facebook“ sraute vartotojas gali apskritai nepastebėti tam tikro įrašo ar nesureaguoti į autoriaus žymą, tuo tarpu tyrimo metu respondentas privalo atkreipti dėmesį į įrašą, į jį įsigilinti ir įvertinti. padėjo jį sumažinti. Tuo tarpu iš anksto atskleistas tyrimo tikslas, gali paskatinti atsakymų šališkumą. Vis dėlto anonimiškumas ir neutralus klausimyno formulavimas tikėtina duomenų tikslumą.
3. Matavimo instrumento apribojimai galėjo paveikti. Kai kurie konstruktai, pavyzdžiui, subjektyvus gebėjimas lengvai atskirti DI turinį nuo žmogaus, ar žinios apie DI įrankius buvo įvertinti vienu teiginiu. Vieno elemento matas turi ribotą patikimumą ir gali neatspindėti visų galimų dimensijos aspektų. Palyginti trumpos skalės buvo naudojamas vertinant ir pasitikėjimą bei požiūrį į DI kaip į įrankį. Reikėtų stengtis, kad vieną matavimo dimensiją sudarytų bent 3 teiginiai. Autentiškumo ir kokybės vertinimas pirminiame tyrimo instrumente buvo sudarytas iš 9 skirtingų skalių. Vis dėlto, po faktorinės analizės šių skalių skaičius sumažėjo iki 3, o tai reiškia kad kokybės ir autentiškumo vertinimui skirti teiginiai buvo pernelyg homogeniški, kad respondentai juos galėtų vertinti skirtingai. Tai galėjo nutikti dėl to, jog tyrimo instrumentas buvo sudarytas remiantis gan dideliu kiekiu skirtingų mokslinių darbų. Vis dėlto, pastarasis sprendimas grindžiamas temos naujumu – išsamių tyrimų, apimančių dirbtinio intelekto ir žmogaus kuriamo turinio įtaką vartotojų pasitikėjimui ir ketinimams kol kas nėra.
4. Tyrime buvo apsiribota dviejų tipų turiniu – tekstiniu socialinio tinklo įrašu su statine nuotrauka. Nors tokio tipo turinys yra bene dažniausiai pasitaikantis socialiniuose tinkluose, egzistuoja ir kiti turinio tipai – vaizdo įrašai, garso turinys ir pan.. Šio tyrimo rezultatai apibendrina tik teksto ir nuotraukos kombinacijos įtaką turinio pasitikėjimo vertinimui ir vartotojų ketinimams. Pagrindinius tyrimo elementus – vaizdą ir tekstą – būtų galima analizuoti individualiai, o po to apjungti į vieną, taip siekiant išsiaiškinti kokią įtaką daro šie elementai individualiai ir palygtini įtaką juos apjungus. Ši prieiga nebuvo pasirinkta, nes tokio tyrimo anketa būtų itin ilga.
5. Eksperimentiam tyrimui buvo pasirinktas racionalaus tipo įrašas. Vertinant kitokio tipo įrašus, pavyzdžiui, transakcinio tipo, rezultatai gali būti kitokie.

Atsižvelgiant į tyrimo diskusiją ir aptartus ribotumus išskiriama keletas **rekomenduotinų tolimesnių tyrimų krypčių**:

1. Rekomenduotina pasitelkti tikimybinę atranką, siekiant sudaryti reprezentatyvią imtį. Taip pat rekomenduotina tyrimu aprėpti ne tik konkrečią šalį, bet ir platesnį kontekstą, pavyzdžiui, Europą. Būtų galima palyginti ir tarpkultūrines tendencijas tarp Lietuvos ir kitų šalių.

2. Tyrimo tikslą ir konkrečius uždavinius atskleisti tik anketos pabaigoje, taip siekiant sumažinti šališkumo tikimybę.
3. Šio tyrimo instrumentą sudarančios skalės buvo sudarytos sintezės būdu, integruojant dešimties šaltinių teiginius. Rekomenduojama konstruoti matavimus remiantis vieno ar kelių autorių sukurtais ir patvirtintais instrumentais, siekiant išvengti homogeniškumo.
4. Rekomenduotina į analizę įtraukti garsinę ar vaizdo įrašų medžiagą, padėsiančią įvertinti platesnį kontekstą, apimančią įvairesnius turinio tipus.
5. Šiam konkrečiam tyrimui buvo pasirinktas racionalaus tipo įrašas socialiniuose tinkluose. Rekomenduotina į tyrimą įtraukti ir emocionalaus bei transakcinio tipo įrašus.
6. Galima būtų išanalizuoti kokią įtaką vartotojų pasitikėjimo pokytis jų ketinimams daro ilgalaikėje perspektyvoje.
7. Rekomenduotina išsiaiškinti kas dar, be pasitikėjimu turiniu, lemia vartotojų ketinimus socialiniuose tinkluose.

Tyrimo išvados ir pasiūlymai

Remiantis tyrimo aktualumo ir problematikos aprašymu, mokslinės literatūros analize ir tyrimo rezultatais, pateikiamos šios baigiamojo magistro darbo išvados:

1. Peržvelgus akademinius šaltinius dirbtinio intelekto ir žmogaus kuriamo vizualiojo ir tekstinio turinio temomis, pastebėtas dviprasmiškas požiūris: DI gali ženkliai pagerinti turinio personalizavimą ir pasiekiamumą, tačiau vartotojų pasitikėjimas tokiu turiniu nėra garantuotas. Pasitikėjimas, tai trąpa, bet tuo pačiu ir esminis sąveikos tarp žmogaus ir DI elementas – jį galima greitai prarasti, jeigu auditorija pasijunta suklaidinta ar neįvertinta. Praktiškai tai reiškia, kad tyrimai apie pasitikėjimą DI kontekste gali padėti gerinti komandų, kuriose yra DI agentai, efektyvumą, sumažinti jų daromas klaidas ir skatinti DI pritaikymą įvairiose organizacijose. Ši spraga pabrėžia poreikį gilintis į šiuos aspektus, siekiant pateikti nuodugnesnę vartotojų elgsenos analizę, vertinant tiek tekstinį, tiek vizualinį turinį. Papildomai svarbu tirti dirbtinio intelekto generuojamo turinio skaidrumą ir paaiškinamumą, kurie daro įtaką vartotojų pasitikėjimui ir norui naudotis tokiomis technologijomis.
2. Dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui ir ketinimams teorinis pagrindimas apima kelis svarbius aspektus, kurie kartu formuoja sudėtingą vartotojų elgsenos socialiniuose tinkluose modelį. Analizės metu paaiškėjo:
 - Egzistuoja labai ryškus skirtumas tarp dirbtinio intelekto ir žmogaus kuriamo turinio, apimantys jų savybes, komunikacines galimybes ir emocinį poveikį vartotojams.
 - DI turinys, nors ir pasižymi efektyvumu, greičiu bei gebėjimu apdoroti didelius duomenų kiekius, dažnai vertinamas kaip mažiau autentiškas dėl savo mechaninio pobūdžio ir riboto gebėjimo perteikti asmeninius išgyvenimus, kūrybiškumą ar kultūrinius niuansus, kuriuos natūraliai įtraukia žmogaus kuriamas turinys. Pastarasis taip pat pasižymi emociniu gilumu, empatija bei personalizacijos galimybėmis.
 - Žmogaus kuriamas turinys dažnai suvokiamas kaip natūralesnis ir nuoširdesnis, nes jis atspindi asmeninę patirtį, emocinį intelektą ir unikalų kūrybinį požiūrį, kuris stiprina vartotojų įsitraukimą ir lojalumą prekės ženklu.
 - DI kuriamas turinys, nors ir gali būti efektyvesnis masinei komunikacijai, dažnai praranda asmeniškumą ir gali būti vertinamas kaip „sintetiškas“ arba „keistas“.
 - Skaidrus turinio kūrėjo identifikavimas padeda mažinti skepticizmą ir kuria pasitikėjimą, nes vartotojai gali lengviau įvertinti šaltinio kompetenciją, patikimumą bei intencijas. Visgi, svarbu pabrėžti, kad autorystės atskleidimas tuo pačiu gali sukelti ir neigiamą reakciją, jei vartotojai skeptiškai vertina DI technologijas, nes DI generuotas turinys gali būti laikomas mažiau patikimu, ypač jei vartotojai suvokia, kad šis turinys nėra natūralus ar neturi emocinio ryšio su auditorija.
 - Literatūrinėje analizėje atkreipiamas dėmesys į socialinių tinklų savitumus, įskaitant platformų funkcionalumą ir vartotojų elgseną. Pavyzdžiui, kaip interaktyvumas, personalizacija ir emocinis poveikis padeda kurti ilgalaikį vartotojų lojalumą, arba kokias galimybes kurti ir skleisti turinio kūrėjams skirtingos socialinių tinklų platformos suteikia.
 - Suformuotas konceptualus modelis išskiria pagrindinius veiksnius, darančius įtaką vartotojų pasitikėjimui turiniu, įskaitant autorystės suvokimą, turinio autentiškumą, kokybę ir vartotojų ketinimus sąveikauti su turiniu, pažymint, kad šie veiksniai tarpusavyje susiję ir gali sustiprinti arba susilpninti vartotojų pasitikėjimą, priklausomai nuo konteksto, autorystės atskleidimo bei vartotojų nuostatų DI atžvilgiu.

3. Remiantis konceptualiuoju modeliu, buvo sudarytas tyrimo modelis ir išskirta 11 hipotezių. Siekiant patikrinti suformuluotas hipotezes, buvo metodiškai apibrėžtas kiekybinis empirinis tyrimas, pasitelkiant eksperimentinius scenarijus ir anketinę apklausą. Tyrimas apėmė dvi respondentų grupes, kurioms pateikti skirtingi turinio scenarijai: viena grupė respondentų vertino DI įrankių pagalba sukurta įrašą, o kita – žmogaus. Rotacija užtikrino tolygų dalyvių pasiskirstymą tarp scenarijų. Tyrimo imtis apėmė pilnamečius Lietuvos gyventojus, turinčius bent vidurinį išsilavinimą, o dalyvių atrankos metodas buvo patogioji atranka internetu, leidusi greitai surinkti reikiamą respondentų skaičių (tikslinis imties dydis ~398, galutinėje apklausoje – 461 dalyvis). Duomenys analizuoti naudojant IBM SPSS programinę įrangą. Ryšiams tarp konstrukto ir kintamųjų nustatyti buvo pasitelkta aprašomoji statistika, faktorinė analizė, koreliacinė ir regresinė analizė, taip pat Mann-Whitney U testas ir chi-kvadrato testas. Tyrimo dalyviams buvo užtikrintas anonimiškumas, savanoriškumo principas ir teisė bet kuriuo metu nutraukti dalyvavimą tyrime.
4. Atliekant tyrimą paaiškėjo, jog respondentų amžius svyravo nuo 18 iki 74 metų, vidutiniškai – 36,25 metų. Didžioji dalis dalyvių (48,8 %) buvo įgiję aukštąjį universitetinį išsilavinimą, o 55,5 % respondentų sudarė moterys. Faktorinės analizės rezultatai lėmė tai, kad pirminės tyrimo skalės buvo itin reikšmingai modifikuotos. Pradinė devynių dimensijų struktūra (sąžiningumas, emocinis ryšys, nuoseklumas, keistumas, sintetiškumas, tikslumas, įtraukimas, aktualumas, suprantamumas) buvo supaprastinta iki šešių dimensijų: „veržlumas ir įtaigumas“, „keistumas“, „sintetiškumas“, „DI kaip bendraautorius“, „DI kaip įrankis“ ir „pasitikėjimas“. Šis pokytis atspindi faktorinės analizės metu identifikuotus konstrukto pergrupavimus, kurie leido tiksliau atspindėti respondentų vertinimus. Atliktas empirinis tyrimas taip pat atskleidė, kad vartotojų suvokiamas turinio veržlumas ir įtaigumas daro teigiamą, o sintetiškumas – reikšmingai neigiamą poveikį pasitikėjimui turiniu. Nustatyta, kad dirbtinio intelekto kurto turinio autoriaus atskleidimas daro neigiamą įtaką pasitikėjimui turiniu, o žmogaus kurto turinio autoriaus atskleidimas – teigiamą. Tuo tarpu keistumas, priešingai nei prognozuota, turėjo teigiamą poveikį pasitikėjimui turiniu, suponuojanti, kad respondentams keistas turinys gali asocijuotis su inovatyvumu ir naujumu. Taip pat nustatyta, kad požiūris į dirbtinį intelektą kaip į bendraautorį moderuoja turinio autoriaus atskleidimo poveikį pasitikėjimui, tačiau požiūris į DI kaip į įrankį šio poveikio reikšmingai neveikia. Tyrimo duomenų analizė taip pat atskleidė, kad vartotojų pasitikėjimo pokytis turiniu teigiamai veikia jų polinkį sąveikauti su turiniu, tačiau šis ryšys buvo gana silpnas, nurodant, kad pasitikėjimas nėra vienintelis veiksnys, lemiantis vartotojų ketinimus. Taip pat pastebėta, kad nepriklausomai nuo respondentų subjektyvių žinių apie DI įrankius ar subjektyvaus gebėjimo atpažinti DI turinį, tinkamai turinio autorių identifikavo 32,5 proc. tyrimo dalyvių. Priešingai nei buvo prognozuota, tyrimo rezultatai atskleidė, kad respondentai, prieš atskleidžiant turinio autorių, DI turinį vertino kaip autentiškesnį, kokybiškesnį ir patikimesnį nei žmogaus sukurta.

Remiantis dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikio vartotojų pasitikėjimui bei ketinimams lyginamosios analizės rezultatais ir teorinėmis implikacijomis šia tema, pateikiamos praktinės **rekomendacijos** organizacijoms ir individualiems asmenims :

1. Tyrimo rezultatai parodė, kad atskleidus turinio autorystę kyla reikšmingi pasitikėjimo pokyčiai – sužinojus, jog įrašą sukūrė DI, vartotojų pasitikėjimas juo smuko, o paaiškėjus, kad autorius žmogus, pasitikėjimas didėjo. Rekomenduojama organizacijoms skaidriai komunikuoti apie turinio autorių, bet tuo pačiu atsižvelgti ir į tai, jog vien tai, kad turinys kaip sukurtas su DI,

pasitikėjimą mažina: šiuo atveju svarbu koreguoti DI turinį ir jį deklaruoti kaip parašytą žmogaus su DI pagalba.

2. Organizacijoms ir socialinių tinklų turinio kūrėjams rekomenduojama kurti turinį su DI pagalba, nes tyrimas atskleidė, kad ir neredaguotas DI turinys gali būti įtaigesnis ir veržlesnis nei žmogaus sukurtas. Tačiau svarbu atkreipti dėmesį į tai, kad šioje vietoje labai svarbu kritiškai vertinti turinio kokybę, pasitelkiant žmogiškąją patirtį.
3. Lietuvoje gyvenantiems socialių tinklų naudotojams, nepriklausomai nuo jų subjektyvių žinių ar subjektyvaus gebėjimo, sunku atskirti DI sukurtą turinį nuo žmogaus. Dėl šios priežasties kyla didelė dezinformacijos rizika. Rekomenduojama organizacijoms ir socialinių tinklų turinio kūrėjams aiškiai ir sąžiningai pažymėti turinį, kuris buvo sukurtas su DI pagalba, taip ribojant dezinformacijos sklaidą.
4. Didesnis pasitikėjimas turiniu statistiškai reikšmingai didina vartotojų ketinimą sąveikauti su turiniu. Dėl šios priežasties organizacijos turėtų užtikrinti, kad jų pateiktas turinys yra tikslus, nuoširdus ir nuoseklus. Svarbu atkreipti dėmesį, kad pasitikėjimas nėra vienintelis faktorius, lemiantis ketinimą sąveikauti, todėl reikia atsižvelgti ir į kitus veiksnius, pavyzdžiui, raginimą sąveikauti.
5. Socialinių tinklų naudotojams rekomenduojama kritiškai vertinti svarbią informaciją socialiniuose tinkluose ir, suabejojus turinio kilmę, pasitelkti kitus DI įrankius tikrojo turinio autoriaus nustatymui. Tyrimo duomenys rodo, kad žmogiškoji patirtis su DI intelekto įrankiais, ar subjektyvi nuomonė apie galėjimą tinkamai identifikuoti turinio autorių, ne visuomet lemia tinkamą turinio autoriaus identifikavimą.
6. Organizacijoms ir socialinių tinklų turinio kūrėjams rekomenduojama neapsiriboti tik tekstinio DI turinio kūrimu, o įtraukti ir vizualinę medžiagą, kurią sukūrė dirbtinis intelektas. Tyrimo rezultatai demonstruoja, kad ir vaizdinę dirbtinio intelekto pagalbą sukurtą medžiagą vartotojai nebūtinai vertins kaip sintetišką. Visgi, šioje vietoje svarbu atkreipti į vizualinės medžiagos kokybę.
7. Tyrimas atskleidė, kad pasitikėjimas paaiškina labai mažą dalį vartotojų ketinimų sąveikauti su turiniu. Tyrėjams rekomenduojama aiškiai, nuosekliai apibrėžti ir operacionalizuoti vartotojų ketinimo sąveikauti su turiniu sąvoką, įtraukiant kiek įmanoma daugiau priežastinių kintamųjų, neapsiribojant tik pasitikėjimu turiniu.
8. Tyrimas analizavo konkrečios šalies tendencijas. Tyrėjams rekomenduojama pakartotinai atlikti tyrimą iš išskirti pastebėjimus, išryškėjusius kitose kultūrose, taip nustatant, ar DI ir žmogaus kuriamo turinio poveikio vartotojų pasitikėjimui bei ketinimams tendencijos yra universalios, ar skiriasi skirtingose auditorijose.

Literatūros sąrašas

1. Afroogh, S., Akbari, A., Malone, E., Kargar, M., & Alambeigi, H. (2024). *Trust in AI: Progress, Challenges, and Future Directions*. <http://arxiv.org/abs/2403.14680>
2. Alkhamees, M., Alsaleem, S., Al-Qurishi, M., Al-Rubaian, M., & Hussain, A. (2021). User trustworthiness in online social networks: A systematic review. *Applied Soft Computing*, 103. <https://doi.org/10.1016/j.asoc.2021.107159>
3. Almira Vania, P., & Daniel Tumpal H., A. (2023). *The Effects Of User-Generated Reviews Versus Influencer-Generated Reviews On Consumer Purchase Intention*. 21(2). <https://doi.org/10.21776/ub.j>
4. Almquist, Y. B., Kvart, S., & Brännström, L. (2019). *A practical guide to quantitative methods with SPSS*. https://su.figshare.com/articles/preprint/A_practical_guide_to_quantitative_methods_with_SPSS/10321829/files/22477367.pdf
5. Altay, S., & Gilardi, F. (2024a). People are skeptical of headlines labeled as AI-generated, even if true or human-made, because they assume full AI automation. *PNAS Nexus*, 3(10). <https://doi.org/10.1093/pnasnexus/pgae403>
6. Belanche, D., Casaló, L. V., & Flavián, M. (2024). Human versus virtual influences, a comparative study. *Journal of Business Research*, 173. <https://doi.org/10.1016/j.jbusres.2023.114493>
7. Berber Sardinha, T. (2024). AI-generated vs human-authored texts: A multidimensional comparison. *Applied Corpus Linguistics*, 4(1). <https://doi.org/10.1016/j.acorp.2023.100083>
8. Briggs, S. R., Cheek, J. M., & Buss, A. H. (1980). An analysis of the Self-Monitoring Scale. *Journal of Personality and Social Psychology*, 38(4), 679–686. <https://doi.org/10.1037/0022-3514.38.4.679>
9. Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., ... Amodei, D. (2020). *Language Models are Few-Shot Learners*. <http://arxiv.org/abs/2005.14165>
10. Buchanan, J., & Hickman, W. (2024). Do people trust humans more than ChatGPT? *Journal of Behavioral and Experimental Economics*, 112. <https://doi.org/10.1016/j.socec.2024.102239>
11. Buzeta, C., De Pelsmacker, P., & Dens, N. (2020). Motivations to Use Different Social Media Types and Their Impact on Consumers' Online Brand-Related Activities (COBRAs). *Journal of Interactive Marketing*, 52, 79–98. <https://doi.org/10.1016/j.intmar.2020.04.004>
12. Campos, D. S., Ferreira de Oliveira, R., de Oliveira Vieira, L., Negreiros de Bragança, P. H., Silva Nunes, J. L., Guimarães, E. C., & Ottoni, F. P. (2023). Revisiting the debate: documenting biodiversity in the age of digital and artificially generated images. *Web Ecology*, 23(2), 135–144. <https://doi.org/10.5194/we-23-135-2023>
13. Chen, X., Xie, H., & Tao, X. (2022). Vision, status, and research topics of Natural Language Processing. *Natural Language Processing Journal*, 1, 100001. <https://doi.org/10.1016/j.nlp.2022.100001>
14. Chen, Y., Prentice, C., Weaven, S., & Hisao, A. (2022). The Influence Of Customer Trust And Artificial Intelligence On Customer Engagement And Loyalty – The Case Of The Home-Sharing Industry. *Frontiers in Psychology*, 13. <https://doi.org/10.3389/fpsyg.2022.912339>

15. Dabbous, A., & Barakat, K. A. (2020). Bridging the online offline gap: Assessing the impact of brands' social network content quality on brand awareness and purchase intention. *Journal of Retailing and Consumer Services*, 53. <https://doi.org/10.1016/j.jretconser.2019.101966>
16. Das, S., Mitra, K., & Mandal, M. (2016). Sample size calculation: Basic principles. In *Indian Journal of Anaesthesia* (Vol. 60, Issue 9, pp. 652–656). Indian Society of Anaesthetists. <https://doi.org/10.4103/0019-5049.190621>
17. Du-Harpur, X., Watt, F. M., Luscombe, N. M., & Lynch, M. D. (2020). What is AI? Applications of artificial intelligence to dermatology. In *British Journal of Dermatology* (Vol. 183, Issue 3, pp. 423–430). Blackwell Publishing Ltd. <https://doi.org/10.1111/bjd.18880>
18. Eigenraam, A. W., Eelen, J., & Verlegh, P. W. J. (2021). Let Me Entertain You? The Importance of Authenticity in Online Customer Engagement. *Journal of Interactive Marketing*, 54, 53–68. <https://doi.org/10.1016/j.intmar.2020.11.001>
19. Felzmann, H., Fosch-Villaronga, E., Lutz, C., & Tamò-Larrieux, A. (2020). Towards Transparency by Design for Artificial Intelligence. *Science and Engineering Ethics*, 26(6), 3333–3361. <https://doi.org/10.1007/s11948-020-00276-4>
20. Ferrario, A., Loi, M., & Viganò, E. (2020). In AI We Trust Incrementally: a Multi-layer Model of Trust to Analyze Human-Artificial Intelligence Interactions. *Philosophy and Technology*, 33(3), 523–539. <https://doi.org/10.1007/s13347-019-00378-3>
21. Geels, J., Graßl, P., Schraffenberger, H., Tanis, M., & Kleemans, M. (2024). Virtual lab coats: The effects of verified source information on social media post credibility. *PLoS ONE*, 19(5 May). <https://doi.org/10.1371/journal.pone.0302323>
22. Gerrard, Y. (2021). What's in a (pseudo)name? Ethical conundrums for the principles of anonymisation in social media research. *Qualitative Research*, 21(5), 686–702. <https://doi.org/10.1177/1468794120922070>
23. Gilbert, J.-P., Ng, V., Niu, J., & Rees, E. E. (2020). A call for an ethical framework when using social media data for artificial intelligence applications in public health research. *Canada Communicable Disease Report*, 169–173. <https://doi.org/10.14745/ccdr.v46i06a03>
24. Goodfellow, I. J., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., & Bengio, Y. (2014). *Generative Adversarial Nets*. <https://doi.org/10.48550/arXiv.1406.2661>
25. Graefe, A., & Bohlken, N. (2020). Automated journalism: A meta-analysis of readers' perceptions of human-written in comparison to automated news. In *Media and Communication* (Vol. 8, Issue 3, pp. 50–59). Cogitatio Press. <https://doi.org/10.17645/mac.v8i3.3019>
26. Gu, C., Jia, S., Lai, J., Chen, R., & Chang, X. (2024). Exploring Consumer Acceptance of AI-Generated Advertisements: From the Perspectives of Perceived Eeriness and Perceived Intelligence. *Journal of Theoretical and Applied Electronic Commerce Research*, 19(3), 2218–2238. <https://doi.org/10.3390/jtaer19030108>
27. Guan, C., Hung, Y. C., & Liu, W. (2022). Cultural differences in hospitality service evaluations: mining insights of user generated content. *Electronic Markets*, 32(3), 1061–1081. <https://doi.org/10.1007/s12525-022-00545-z>
28. Hochstein, R. E., Harmeling, C. M., & Perko, T. (2023). Toward a theory of consumer digital trust: Meta-analytic evidence of its role in the effectiveness of user-generated content. *Journal of the Academy of Marketing Science*. <https://doi.org/10.1007/s11747-023-00982-y>
29. Hofacker, C., Golgeci, I., Pillai, K. G., & Gligor, D. M. (2020). Digital marketing and business-to-business relationships: a close look at the interface and a roadmap for the future. In *European*

- Journal of Marketing* (Vol. 54, Issue 6, pp. 1161–1179). Emerald Group Holdings Ltd. <https://doi.org/10.1108/EJM-04-2020-0247>
30. Hofeditz, L., Mirbabaie, M., Holstein, J., & Stieglitz, S. (2021) Do you Trust an AI-Journalist? A Credibility Analysis of News Content with AI-Authorship. https://aisel.aisnet.org/ecis2021_rp/50
 31. Holone, H. (2016). The filter bubble and its effect on online personal health information. *Croatian Medical Journal*, 57(3), 298–301. <https://doi.org/10.3325/cmj.2016.57.298>
 32. Huschens, M., Briesch, M., Sobania, D., & Rothlauf, F. (2023). *Do You Trust ChatGPT? -- Perceived Credibility of Human and AI-Generated Content*. <http://arxiv.org/abs/2309.02524>
 33. Ismagilova, E., Slade, E., Rana, N. P., & Dwivedi, Y. K. (2020). The effect of characteristics of source credibility on consumer behaviour: A meta-analysis. *Journal of Retailing and Consumer Services*, 53. <https://doi.org/10.1016/j.jretconser.2019.01.005>
 34. Ivanov, S. (2023). *Using Artificial Intelligence to Create Marketing Content-Opportunities and Limitations*. <https://www.researchgate.net/publication/372187456>
 35. Jamieson, M. K., Govaart, G. H., & Pownall, M. (2023). Reflexivity in quantitative research: A rationale and beginner's guide. In *Social and Personality Psychology Compass* (Vol. 17, Issue 4). John Wiley and Sons Inc. <https://doi.org/10.1111/spc3.12735>
 36. Jenkins, E. L., Ilcic, J., Barklamb, A. M., & McCaffrey, T. A. (2020b). Assessing the credibility and authenticity of social media content for applications in health communication: Scoping review. In *Journal of Medical Internet Research* (Vol. 22, Issue 7). JMIR Publications Inc. <https://doi.org/10.2196/17296>
 37. Ji, G. M., Cheah, J. H., Sigala, M., Ng, S. I., & Choo, W. C. (2023). Tell me about your culture, to predict your tourism activity preferences and evaluations: cross-country evidence based on user-generated content. *Asia Pacific Journal of Tourism Research*, 28(10), 1052–1070. <https://doi.org/10.1080/10941665.2023.2283599>
 38. Joy, S., Saranya, C., Sinthu, R., & Swetha, M. (2024). Unpacking The Ethical Implications Of Ai In Social Media. *International Research Journal of Modernization in Engineering Technology and Science*, 06(11). https://www.irjmets.com/uploadedfiles/paper//issue_11_november_2024/63960/final/fin_irjmets1731821089.pdf
 39. Kafle, S. C. (2023). *Correlation and Regression Analysis Using SPSS*. <https://journal.oxfordcollege.edu.np/index.php/ojmts/article/view/14>
 40. Kelly, S., Kaye, S. A., & Oviedo-Trespalacios, O. (2023). What factors contribute to the acceptance of artificial intelligence? A systematic review. *Telematics and Informatics*, 77. <https://doi.org/10.1016/j.tele.2022.101925>
 41. Kemp, E., Porter, M., Anaza, N. A., & Min, D. J. (2021). The impact of storytelling in creating firm and customer connections in online environments. *Journal of Research in Interactive Marketing*, 15(1), 104–124. <https://doi.org/10.1108/JRIM-06-2020-0136>
 42. Kim, J. Y., Kioussis, S., & Molleda, J. C. (2015). Use of affect in blog communication: Trust, Credibility, and Authenticity. *Public Relations Review*, 41(4), 504–507. <https://doi.org/10.1016/j.pubrev.2015.07.002>
 43. Krishen, A. S., Dwivedi, Y. K., Bindu, N., & Kumar, K. S. (2021). A broad overview of interactive digital marketing: A bibliometric network analysis. In *Journal of Business Research* (Vol. 131, pp. 183–195). Elsevier Inc. <https://doi.org/10.1016/j.jbusres.2021.03.061>

44. Lappeman, J., Marlie, S., Johnson, T., & Poggenpoel, S. (2023). Trust and digital privacy: willingness to disclose personal information to banking chatbot services. *Journal of Financial Services Marketing*, 28(2), 337–357. <https://doi.org/10.1057/s41264-022-00154-z>
45. Larsson, S., & Heintz, F. (2020). Transparency in artificial intelligence. *Internet Policy Review*, 9(2), 1–16. <https://doi.org/10.14763/2020.2.1469>
46. Lee, E. J. (2020). Authenticity Model of (Mass-Oriented) Computer-Mediated Communication: Conceptual Explorations and Testable Propositions. *Journal of Computer-Mediated Communication*, 25(1), 60–73. <https://doi.org/10.1093/jcmc/zmz025>
47. Li, F., & Yang, Y. (2024). Impact of Artificial Intelligence-Generated Content Labels On Perceived Accuracy, Message Credibility, and Sharing Intentions for Misinformation: Web-Based, Randomized, Controlled Experiment. *JMIR Formative Research*, 8. <https://doi.org/10.2196/60024>
48. Lim, W. M., & Rasul, T. (2022). Customer engagement and social media: Revisiting the past to inform the future. *Journal of Business Research*, 148, 325–342. <https://doi.org/10.1016/j.jbusres.2022.04.068>
49. Loconsole, C., & Panarari, M. (2025). *Algorithmic Journalism and Ideological Polarization: An Experimental Work Around ChatGPT and the Production of Politically Oriented Information*. <https://doi.org/10.54941/ahfe1005943>
50. Maares, P., Banjac, S., & Hanusch, F. (2021). The labour of visual authenticity on social media: Exploring producers' and audiences' perceptions on Instagram. *Poetics*, 84. <https://doi.org/10.1016/j.poetic.2020.101502>
51. McDermid, J. A., Jia, Y., Porter, Z., & Habli, I. (2021). Artificial intelligence explainability: The technical and ethical dimensions. In *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences* (Vol. 379, Issue 2207). Royal Society Publishing. <https://doi.org/10.1098/rsta.2020.0363>
52. Mohamed, E. A. S., Osman, M. E., & Mohamed, B. A. (2024a). The Impact of Artificial Intelligence on Social Media Content. *Journal of Social Sciences*, 20(1), 12–16. <https://doi.org/10.3844/jssp.2024.12.16>
53. Mohamed, N. N., Jaafar, N., & Ayupp, K. (2023a). The Influence of User-Generated Content Information Credibility and Information Adoption on Consumer Purchase Intention. *International Journal of Academic Research in Business and Social Sciences*, 13(2). <https://doi.org/10.6007/ijarbss/v13-i2/16269>
54. Mohamed, N. N., Jaafar, N., & Ayupp, K. (2023b). *The Mediating Effect of Information Adoption on The Association between Social Media Influencer Information Credibility and Purchase Intention*.
55. Monteith, S., Glenn, T., Geddes, J. R., Whybrow, P. C., Achtyes, E., & Bauer, M. (2024). Artificial intelligence and increasing misinformation. In *British Journal of Psychiatry* (Vol. 224, Issue 2, pp. 33–35). Cambridge University Press. <https://doi.org/10.1192/bjp.2023.136>
56. Niu Yongbo. (2025). From Human Anchors to AI Anchors: A Review of Technology, Ethics, and Audience Response in Audiovisual Media Transformation. *International Theory and Practice in Humanities and Social Sciences*, 2(2), 361–372. <https://doi.org/10.70693/itphss.v2i2.166>
57. Novelli, C., Taddeo, M., & Floridi, L. (2024). Accountability in artificial intelligence: what it is and how it works. *AI and Society*, 39(4), 1871–1882. <https://doi.org/10.1007/s00146-023-01635-y>

58. Ostic, D., Qalati, S. A., Barbosa, B., Shah, S. M. M., Galvan Vela, E., Herzallah, A. M., & Liu, F. (2021). Effects of Social Media Use on Psychological Well-Being: A Mediated Model. *Frontiers in Psychology*, 12. <https://doi.org/10.3389/fpsyg.2021.678766>
59. Park, J., Oh, C., & Kim, H. Y. (2024). AI vs. human-generated content and accounts on Instagram: User preferences, evaluations, and ethical considerations. *Technology in Society*, 79. <https://doi.org/10.1016/j.techsoc.2024.102705>
60. Piduru, B. R. (2023a). The Role of Artificial Intelligence in Content Personalization: Transforming User Experience in the Digital Age. *Journal of Artificial Intelligence & Cloud Computing*, 1–5. [https://doi.org/10.47363/JAICC/2023\(2\)193](https://doi.org/10.47363/JAICC/2023(2)193)
61. Piligrimienė, Ž. (2016). Marketingo tyrimų duomenų analizė SPSS programa. Kaunas: Technologija, 124 p., ISBN 978-609-02-1219-6
62. Pustahelyi, R. (2020). *Emotional Ai And Its Challenges In The Viewpoint Of Online Marketing*. DOI: <https://www.researchgate.net/publication/343017664>
63. Qi, W., Pan, J., Lyu, H., & Luo, J. (2024). Excitements and concerns in the post-ChatGPT era: Deciphering public perception of AI through social media analysis. *Telematics and Informatics*, 92. <https://doi.org/10.1016/j.tele.2024.102158>
64. Radford, A., Kim, J. W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., Krueger, G., & Sutskever, I. (2021). *Learning Transferable Visual Models From Natural Language Supervision*. <http://arxiv.org/abs/2103.00020>
65. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., & Ommer, B. (2021). *High-Resolution Image Synthesis with Latent Diffusion Models*. <http://arxiv.org/abs/2112.10752>
66. Saheb, T., Sidaoui, M., & Schmarzo, B. (2024). Convergence of artificial intelligence with social media: A bibliometric & qualitative analysis. In *Telematics and Informatics Reports* (Vol. 14). Elsevier B.V. <https://doi.org/10.1016/j.teler.2024.100146>
67. Sardá, T., Natale, S., Sotirakopoulos, N., & Monaghan, M. (2019). Understanding online anonymity. *Media, Culture and Society*, 41(4), 557–564. <https://doi.org/10.1177/0163443719842074>
68. Sargin, S. (2024). Antecedents and Consequences of Consumers Attitudes Towards Artificial Intelligence in Social Media. *Business and Economics Research Journal*. <https://doi.org/10.20409/berj.2024.443>
69. Scharowski, N., Perrig, S. A. C., Svab, M., Opwis, K., & Brühlmann, F. (2023). Exploring the effects of human-centered AI explanations on trust and reliance. *Frontiers in Computer Science*, 5. <https://doi.org/10.3389/fcomp.2023.1151150>
70. Schlosser, A. E. (2020). Self-disclosure versus self-presentation on social media. In *Current Opinion in Psychology* (Vol. 31, pp. 1–6). Elsevier B.V. <https://doi.org/10.1016/j.copsyc.2019.06.025>
71. Schreiner, M., Fischer, T., & Riedl, R. (2021). Impact of content characteristics and emotion on behavioral engagement in social media: literature review and research agenda. *Electronic Commerce Research*, 21(2), 329–345. <https://doi.org/10.1007/s10660-019-09353-8>
72. Seiler, R., & Kucza, G. (2017). Source Credibility Model, Source Attractiveness Model And Match-Up-Hypothesis-An Integrated Model. In *Economy & Business ISSN* (Vol. 11). https://www.researchgate.net/publication/319448379_Source_Credibility_Model_Source_Attractiveness_Model_And_Match-Up-Hypothesis-An_Integrated_Model
73. Shahbaznezhad, H., Dolan, R., & Rashidirad, M. (2021). The Role of Social Media Content Format and Platform in Users' Engagement Behavior. *Journal of Interactive Marketing*, 53, 47–65. <https://doi.org/10.1016/j.intmar.2020.05.001>

74. Singhal, A., Neveditsin, N., Tanveer, H., & Mago, V. (2024a). Toward Fairness, Accountability, Transparency, and Ethics in AI for Social Media and Health Care: Scoping Review. In *JMIR Medical Informatics* (Vol. 12). JMIR Publications Inc. <https://doi.org/10.2196/50048>
75. Smith, D., Matthews, J. (2021). Johnson And Journalism: Anonymous Sources In Senior Journalists' Social Media Feeds. *Populism The Pandemic and The Media: Journalism in the age of Brexit, Covid ,Donald Trump and Boris Johnson* (pp.33-38). *Prieiga per internetą*: <https://www.researchgate.net/publication/352019834>
76. Stahl, B. C., Antoniou, J., Ryan, M., Macnish, K., & Jiya, T. (2021). Organisational responses to the ethical issues of artificial intelligence. *AI and Society*, 37(1), 23–37. <https://doi.org/10.1007/s00146-021-01148-6>
77. Taber, K. S. (2018). The Use of Cronbach's Alpha When Developing and Reporting Research Instruments in Science Education. *Research in Science Education*, 48(6), 1273–1296. <https://doi.org/10.1007/s11165-016-9602-2>
78. Ulfert, A. S., Georganta, E., Centeio Jorge, C., Mehrotra, S., & Tielman, M. (2024). Shaping a multidisciplinary understanding of team trust in human-AI teams: a theoretical framework. *European Journal of Work and Organizational Psychology*, 33(2), 158–171. <https://doi.org/10.1080/1359432X.2023.2200172>
79. Wang, H. (2022). Transparency as Manipulation? Uncovering the Disciplinary Power of Algorithmic Transparency. *Philosophy and Technology*, 35(3). <https://doi.org/10.1007/s13347-022-00564-w>
80. Wei, S., Xu, X., Qi, X., Yin, X., Xia, J., Ren, J., Tang, P., Zhong, Y., Chen, Y., Ren, X., Liang, Y., Huang, L., Xie, K., Gui, W., Tan, W., Sun, S., Hu, Y., Liu, Q., Li, N., ... Xie, Y. (2023). *AcademicGPT: Empowering Academic Research*. <http://arxiv.org/abs/2311.12315>
81. Xie, Y., Rawal, A., Cen, Y., Zhao, D., Narang, S. K., & Sushmita, S. (2024). *MUGC: Machine Generated versus User Generated Content Detection*. <http://arxiv.org/abs/2403.19725>
82. Yang, J., Teran, C., Battocchio, A. F., Bertellotti, E., & Wrzesinski, S. (2021). Building Brand Authenticity on Social Media: The Impact of Instagram Ad Model Genuineness and Trustworthiness on Perceived Brand Authenticity and Consumer Responses. *Journal of Interactive Advertising*, 21(1), 34–48. <https://doi.org/10.1080/15252019.2020.1860168>
83. Yang, M., & Rau, P. L. P. (2024). Trust building with artificial intelligence: comparing with human in investment behaviour, emotional arousal and neuro activities. *Theoretical Issues in Ergonomics Science*, 25(5), 593–614. <https://doi.org/10.1080/1463922X.2023.2290019>
84. Yaqub, W., Kakhidze, O., Brockman, M. L., Memon, N., & Patil, S. (2020, April 21). Effects of Credibility Indicators on Social Media News Sharing Intent. *Conference on Human Factors in Computing Systems - Proceedings*. <https://doi.org/10.1145/3313831.3376213>
85. Yeung, A., Ng, E., & Abi-Jaoude, E. (2022). TikTok and Attention-Deficit/Hyperactivity Disorder: A Cross-Sectional Study of Social Media Content Quality. *Canadian Journal of Psychiatry*, 67(12), 899–906. <https://doi.org/10.1177/07067437221082854>
86. Yildiz Durak, H., Eğin, F., & Onan, A. (2025). A Comparison of Human-Written Versus AI-Generated Text in Discussions at Educational Settings: Investigating Features for ChatGPT, Gemini and BingAI. *European Journal of Education*, 60(1). <https://doi.org/10.1111/ejed.70014>
87. Zhang, Y., Gaggiano, J. D., Yongsatianchot, N., Suhaimi, N. M., Kim, M., Sun, Y., Griffin, J., & Parker, A. G. (2023, April 19). What Do We Mean When We Talk about Trust in Social Media? A Systematic Review. *Conference on Human Factors in Computing Systems - Proceedings*. <https://doi.org/10.1145/3544548.3581019>

Informacijos šaltinių sąrašas

1. Always. Always #LikeAGirl [žiūrėta 2025-01-07]. Prieiga per internetą: <https://www.always.com/en-us/about-us/our-epic-battle-like-a-girl>
2. Andrew Hayes. The PROCESS macro for SPSS, SAS, and R [žiūrėta 2025-05-01]. Prieiga per internetą: <https://www.processmacro.org/index.html>
3. PlannerBear. Automate ir Scale Your Marketing [žiūrėta 2025-02-20]. Prieiga per internetą: <https://www.bannerbear.com/>
4. Canva. What will you design today? [žiūrėta 2025-02-20]. Prieiga per internetą: <https://www.canva.com/>
5. ChatDesk. #InMyDenim Guess [žiūrėta 2025-01-07]. Prieiga per internetą: <https://www.chatdesk.com/tiktok-campaigns-and-case-studies/inmydenim-guess-hashtag-challenge>
6. Deloitte (2024). Earning trust as gen AI takes hold: 2024 Connected Consumer Survey [žiūrėta 2025-02-28]. Prieiga per internetą: <https://www2.deloitte.com/us/en/insights/industry/telecommunications/connectivity-mobile-trends-survey.html>
7. Dove. Campaigns [žiūrėta 2025-01-07]. Prieiga per internetą: <https://www.dove.com/uk/stories/campaigns.html>
8. European Union (2024). AI Act [žiūrėta 2025-02-27]. Prieiga per internetą: <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>
9. European Union (2016). General Data Protection Regulation [žiūrėta 2025-04-01]. Prieiga per internetą: <https://eur-lex.europa.eu/eli/reg/2016/679/oj>
10. IBM (2025a). IBM Study: AI Spending Expected to Surge 52% Beyond IT Budgets as Retail Brands Embrace Enterprise-Wide Innovation [žiūrėta 2025-01-10]. Prieiga per internetą: <https://newsroom.ibm.com/2025-01-07-ibm-study-ai-spending-expected-to-surge-52-beyond-it-budgets-as-retail-brands-embrace-enterprise-wide-innovation>
11. IBM (2025b). Prompt engineering guide [žiūrėta 2025-04-21]. Prieiga per internetą: <https://www.ibm.com/think/topics/prompt-engineering-guide>
12. IBM (2023). AI in marketing: How to leverage this powerful new technology for your next campaign? [žiūrėta 2025-01-10]. Prieiga per internetą: <https://www.ibm.com/think/topics/ai-in-marketing>
13. Justicija.eu. Kokie duomenys gali būti laikomi asmens duomenimis ir kaip tinkamai jie turi būti tvarkomi? [žiūrėta 2025-04-10]. Prieiga per internetą: <https://justicija.eu/kokie-duomenys-gali-buti-laikomi-asmens-duomenimis-ir-kaip-tinkamai-jie-turi-buti-tvarkomi/>
14. Lego. The world is playable, so let's play [žiūrėta 2025-01-07]. Prieiga per internetą: <https://www.lego.com/lt-lt/rebuild-the-world?page=2>
15. McDonald's blog. The Wholesome McAlloo Tikki Burger [žiūrėta 2025-01-07]. Prieiga per internetą: <https://mcdonaldsblog.in/2018/05/the-wholesome-mcalloo-tikki-burger/>
16. Nike. Great Athletes Remind the World There's Nothing Wrong With Wanting to Win [žiūrėta 2025-01-07]. Prieiga per internetą: <https://about.nike.com/en/newsroom/releases/winning-isnt-for-everyone-campaign>
17. Nosto (2021). Post-Pandemic Shifts in Consumer Shopping Habits: Authenticity, Personalization and the Power of UGC [žiūrėta 2025-02-28]. Prieiga per internetą: <http://nosto.com/wp-content/uploads/2021/08/Stackla-Post-Pandemic-Shifts-in-Consumer-Shopping-Habits-Data-Report FINAL compressed.pdf>

18. Oficialiosios statistikos portalas (2025). Duomenų bazė [žiūrėta 2025-04-01]. Prieiga per internetą: <https://osp.stat.gov.lt/statistiniu-rodikliu-analize#/>
19. Oficialiosios statistikos portalas (2025). Tautybė, gimtoji kalba ir tikyba [žiūrėta 2025-04-01]. Prieiga per internetą: <https://osp.stat.gov.lt/2021-gyventoju-ir-bustu-surasy-mo-rezultatai/tautybe-gimtoji-kalba-ir-tikyba>
20. Bitės profai (2024). Socialinių tinklų naudojimas Lietuvoje [žiūrėta 2025-04-01]. Prieiga per internetą: <https://www.bite.lt/sites/default/files/socialiniu-tinklu-naudojimas-lietuvoje.pdf>
21. PennState (2022). Users trust AI as much as humans for flagging problematic content [žiūrėta 2025-02-28]. Prieiga per internetą: <https://www.psu.edu/news/institute-computational-and-data-sciences/story/users-trust-ai-much-humans-flagging-problematic>
22. Sodra (2023). Gyventojų darbo pajamų apžvalga [žiūrėta 2025-03-26]. Prieiga per internetą: https://www.sodra.lt/uploads/documents/files/20230831_Sodra_2023%20%20ketvircio%20darbo%20pajam%C5%B3%20ap%C5%BEvalga.pdf
23. Statista. Most popular social media websites in Lithuania in 2023, based on share of referred web traffic on third-party websites. [žiūrėta 2025-01-07]. Prieiga per internetą: <https://www.statista.com/statistics/1165917/market-share-of-the-most-popular-social-media-websites-in-lithuania/>
24. Spotify. 2024 Wrapped Is Here! Experience Your Year in Listening Like Never Before [žiūrėta 2025-01-07]. Prieiga per internetą: <https://newsroom.spotify.com/2024-12-04/wrapped-user-experience-2024/>
25. The University of Kansas. Study finds readers trust news less when AI is involved, even when they don't understand to what extent [žiūrėta 2025-02-28]. Prieiga per internetą: <https://news.ku.edu/news/article/study-finds-readers-trust-news-less-when-ai-is-involved-even-when-they-dont-understand-to-what-extent>
26. Valstybinė duomenų apsaugos inspekcija. Bendrovei pareikštas papeikimas ir skirti nurodymai dėl neteisėto vaikų atvaizdų viešinimo per „Google Photos“ ir teisės susipažinti su asmens duomenimis pažeidimų [žiūrėta 2025-04-10]. Prieiga per internetą: https://vdai.lrv.lt/public/canonical/1734939897/683/4_2024%20Sprendimo%20apibendrinimas%20BDAR%205%20str%201dalies%20a%20ir%2015%20str.pdf

Priedai

1 priedas. Pirminės tyrimo hipotezės ir jų detalizacija

Nr.	Hipotezė	Detalizavimas
H1	Vartotojų suvokiamas turinio autentiškumas teigiamai veikia jų pasitikėjimą turiniu.	Jei turinys atrodo sąžiningas, emocionalus, nuoseklus, nesukelia keistumo jausmo ir nėra pernelyg sintetiškas, jis vartotojui atrodo autentiškas, o tai didina pasitikėjimą (Almira Vania & Daniel Tumpal H., 2023; Belanche et al., 2024; Buzeta et al., 2020; Gu et al., 2024).
H1a	Vartotojų suvokiamas turinio sąžiningumas daro teigiamą įtaką turinio autentiškumo vertinimui	Sąžiningas turinys vartotojams atrodo tikroviškesnis, todėl skatina pasitikėjimą turiniu (Almira Vania & Daniel Tumpal H., 2023; Belanche et al., 2024).
H1b	Vartotojų suvokiamas turinio emocinis ryšys daro teigiamą įtaką turinio autentiškumo vertinimui	Turinys, kuriame kuriamas emocinis ryšys, vartotojui atrodo autentiškesnis ir kelia daugiau pasitikėjimo (Almira Vania & Daniel Tumpal H., 2023; Buzeta et al., 2020; Gu et al., 2024).
H1c	Vartotojų suvokiamas turinio nuoseklumas daro teigiamą įtaką turinio autentiškumo vertinimui	Nuoseklus turinys, be prieštaravimų ir nenuoseklumų, stiprina vartotojų suvokimą apie autentiškumą (Gu et al., 2024; Sargin, 2024).
H1d	Vartotojų suvokiamas turinio keistumas daro neigiamą įtaką turinio autentiškumo vertinimui	Turinys, sukeliantis keistumo ar nejaukumo pojūtį, mažina autentiškumo suvokimą ir pasitikėjimą (Gu et al., 2024).
H1e	Vartotojų suvokiamas turinio sintetiškumas daro neigiamą įtaką turinio autentiškumo vertinimui	Jei turinys atrodo pernelyg sintetiškas ar „dirbtinis“, vartotojai jį vertina kaip mažiau autentišką ir menkliau patikimą (Gu et al., 2024).
H2	Vartotojų suvokiama turinio kokybė teigiamai veikia jų pasitikėjimą turiniu.	Tikslus, įtraukiantis, aktualus, aiškus ir inovatyvus turinys vartotojui atrodo kokybiškas, todėl skatina jį labiau pasitikėti matoma informacija (Almira Vania & Daniel Tumpal H., 2023; Belanche et al., 2024; Gu et al., 2024; Huschens et al., 2023).
H2a	Vartotojų suvokiamas turinio tikslumas daro teigiamą įtaką turinio kokybės vertinimui	Tikslus turinys, besiremiantis faktais ir duomenimis, formuoja vartotojo įsitikinimą apie turinio patikimumą (Huschens et al., 2023)
H2b	Vartotojų suvokiamas turinio gebėjimas įtraukti daro teigiamą įtaką turinio kokybės vertinimui	Įtraukiantis turinys patraukia dėmesį ir palaiko vartotojų įsitraukimą, stiprindamas kokybės išspūdį (Huschens et al., 2023).
H2c	Vartotojų suvokiamas turinio aktualumas daro teigiamą įtaką turinio kokybės vertinimui	Turinys, kuris yra aktualus vartotojui, suvokiamas kaip vertingesnis ir patikimesnis (Belanche et al., 2024)
H2d	Vartotojų suvokiamas turinio suprantamumas daro teigiamą įtaką turinio kokybės vertinimui	Aiškiai ir suprantamai pateiktas turinys vartotojui atrodo kokybiškesnis (Huschens et al., 2023).
H3	Tikrojo turinio autoriaus atskleidimas moderuoja pasitikėjimą turiniu.	Atvirai atskleidus, kad turinį sukūrė dirbtinis intelektas, vartotojų pasitikėjimas sumažėja dėl išankstinių nuostatų ir skeptiškumo, tuo tarpu žmogaus kurtu turiniu vartotojai yra labiau linkę pasitikėti, dėl suvokiamo tikrumo ir emocionalumo (Altay & Gilardi, 2024b). Kita vertus, kituose moksliniuose darbuose pastebima, kad kai žmonėms aiškiai pasakoma, kas parašė tekstą, pasitikėjimas būna vienodai žemas tiek DI, tiek žmogaus tekstams (Buchanan & Hickman, 2024)
H3a	Dirbtinio intelekto kurto turinio autoriaus atskleidimas daro neigiamą įtaką pasitikėjimui turiniu	
H3b	Žmogaus kurto turinio autoriaus atskleidimas daro teigiamą įtaką pasitikėjimui turiniu	

		Tai rodo, kad autorystės atskleidimas skatina kritinį vertinimą nepriklausomai nuo autoriaus tipo.
H4	Vartotojų pasitikėjimas turiniu teigiamai veikia jų polinkį sąveikauti su turiniu.	Kuo labiau vartotojas pasitiki turiniu, tuo didesnė tikimybė, kad jis sąveikaus su turiniu – komentuodamas, dalindamasis ar reaguodamas (Dabbous & Barakat, 2020; Gu et al., 2024; Park et al., 2024).
H5	Požiūris į DI moderoja kaip vartotojų pasitikėjimą veikia tikrojo turinio autoriaus atskleidimas.	Požiūris į DI – ar jis laikomas įrankiu, ar bendraautoriumi – daro reikšmingą poveikį tam, kaip vartotojas vertina atskleistą autorystę. Tyrimų rezultatai rodo, kad respondentai labiau linkę žiūrėti į DI kaip į įrankį, o ne kaip į kūrėją, o tai keičia jų pasitikėjimą turiniu: vartotojai, laikantys DI bendraautoriumi, dažniau jį toleruoja, vertina skaidrų autorystės atskleidimą, tuo tarpu DI kaip įrankį vertinantys asmenys linkę mažiau pasitikėti turiniu dėl nuoširdumo trūkumo (Gu et al., 2024; E. A. S. Mohamed et al., 2024a).
H5a	Dirbtinio intelekto turinio tikrojo autoriaus atskleidimas teigiamai veikia vartotojų, kurie į DI žvelgia kaip į bendraautorį, pasitikėjimą turiniu.	
H5b	Dirbtinio intelekto turinio tikrojo autoriaus atskleidimas neigiamai veikia vartotojų, kurie į DI žvelgia kaip į įrankį, pasitikėjimą turiniu.	
H6	Vartotojų žinios apie DI įrankius teigiamai veikia jų gebėjimą identifikuoti tikrąjį turinio autorių	Kuo daugiau vartotojas žino apie DI įrankius, tuo didesnė tikimybė, kad jis gebės atpažinti ar bent jau įtarti turinio kilmę (Gu et al., 2024; Xie et al., 2024). Visgi Park et al. (2024) pastebi, kad dabartinis DI turinys yra toks įtikinamas, kad žmogui atskirti kas yra tikrasis turinio autorius yra labai sunku.
H7	Vartotojų suvokiamas gebėjimas tinkamai identifikuoti DI turinį teigiamai veikia jų gebėjimą identifikuoti DI turinį.	Pasitikėjimas savo sugebėjimu atskirti žmogaus ir DI sukurtą turinį lemia didesnę identifikavimo tikslumą – subjektyvus tikėjimas savo kompetencija atskleidžia gebėjimą atkreipti dėmesį į tinkamas detales (Xie et al., 2024).
H8	Vartotojų suvokiamas pasitikėjimo turiniu didėjimas žinant autorių teigiamai veikia jų pasitikėjimą turiniu, atskleidus tikrąjį turinio autorių.	Vartotojai vertina turinio skaidrumą ir yra linkę pasitikėti turiniu labiau, kai žino jo autorių (Gu et al., 2024).
H9	Žmogaus sukurtas turinys daro stipresnį teigiamą poveikį vartotojų suvokiamam turinio autentiškumui nei dirbtinio intelekto sukurtas turinys	Žmogaus kuriamas turinys pasižymi kultūriniu kontekstu, spontaniškumu ir vertybių raiška – tai stiprina jo autentiškumo suvokimą (Almira Vania & Daniel Tumpal H., 2023; Belanche et al., 2024; Buzeta et al., 2020; Gu et al., 2024). Autentiškesniu laikomas turinys, kuris yra sąžiningesnis, emociškai paveikesnis, nuoseklesnis, mažiau keistas ir sintetiškas (Almira Vania & Daniel Tumpal H., 2023; Belanche et al., 2024; Buzeta et al., 2020; Gu et al., 2024)
H9a	Žmogaus sukurtas turinys daro stipresnį teigiamą poveikį vartotojų suvokiamam turinio sąžiningumui nei dirbtinio intelekto sukurtas turinys	
H9b	Žmogaus sukurtas turinys daro stipresnį teigiamą poveikį vartotojų suvokiamam turinio emociniam ryšiui nei dirbtinio intelekto sukurtas turinys	
H9c	Žmogaus sukurtas turinys daro stipresnį teigiamą poveikį vartotojų suvokiamam turinio nuoseklumui nei dirbtinio intelekto sukurtas turinys	
H9d	Žmogaus sukurtas turinys daro stipresnį neigiamą poveikį vartotojų suvokiamam turinio keistumui nei dirbtinio intelekto sukurtas turinys	
H9e	Žmogaus sukurtas turinys daro stipresnį neigiamą poveikį vartotojų suvokiamam turinio sintetiškumui nei dirbtinio intelekto sukurtas turinys	
H10	Žmogaus sukurtas turinys daro stipresnį teigiamą poveikį vartotojų suvokiamam turinio kokybei nei dirbtinio intelekto sukurtas turinys	Žmogaus sukurtas turinys turi stipresnį „kokybės patikimumo“ efektą vartotojų akyse dėl jo giluminės prasmės, vertybinio pagrindo ir natūralaus stilistinio

H10a	Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio tikslumui nei dirbtinio intelekto sukurtas turinys	įvairovės, kurių DI dažnai nepajėgia atkurti (Berber Sardinha, 2024; Dabbous & Barakat, 2020; Huschens et al., 2023; Kemp et al., 2021).
H10b	Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio gebėjimui įtraukti nei dirbtinio intelekto sukurtas turinys	
H10c	Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio aktualumui nei dirbtinio intelekto sukurtas turinys	
H10d	Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų suvokiamam turinio suprantamumui nei dirbtinio intelekto sukurtas turinys	
H11	Žmogaus sukurtas turinys daro stipresnę teigiamą poveikį vartotojų pasitikėjimui turiniu nei dirbtinio intelekto sukurtas turinys.	Žmonės, manydami, jog turinį parengė žmogus, yra linkę labiau juo pasitikėti. Pastebimas išankstinis šališkumas žmogaus naudai, kai respondentas nežino autoriaus, bet remiasi įtarimais ar stiliumi (Buchanan & Hickman, 2024).

2 priedas. Tyrimo anketa

Žemiau pateikta tyrimo anketa. Tekstas parašytas *pakreiptu šriftu* anketoje nebuvo matomas. Į visus klausimus respondentai turėjo atsakyti pasirinkdami vieną variantą, nebent nurodyta kitaip.

Įvadas. Sveiki, Esu Airidas Kundrotas, Kauno technologijos universiteto Marketingo valdymo magistro studijų studentas. Šiuo metu atlieku baigiamojo magistro darbo projekto tyrimą, kurio tikslas – palyginti dirbtinio intelekto ir žmogaus kuriamo turinio socialiniuose tinkluose poveikį vartotojų pasitikėjimui bei ketinimams. Tyrimo rezultatai padės geriau suprasti, kaip žmonės vertina žmogaus ir dirbtinio intelekto kuriamo turinio kokybę ir autentiškumą, kokią įtaką šis vertinimas daro pasitikėjimui turiniu ir kokius veiksmus skatina. Jūsų dalyvavimas apklausoje yra visiškai savanoriškas ir anonimiškas, o visi pateikti duomenys – konfidencialūs. Surinkta informacija bus naudojama apibendrintai ir tik baigiamojo magistro rengimo tikslais. Anketos pildymas truks apie 10–12 minučių. Jei turėtumėte klausimų arba norėtumėte susipažinti su tyrimo rezultatais, kviečiu susisiekti el. paštu airidas.kundrotas@ktu.edu Dėkoju už Jūsų laiką ir dalyvavimą! 1 klausimų blokas. Kaip dažnai naudojate socialiniais tinklais? • Kartą per dieną • Keletą kartų per savaitę • Kartą per savaitę • Keletą kartų per mėnesį • Kartą per mėnesį • Rečiau • Niekada Kiek laiko praleidžiate naudodamiesi socialiniais tinklais per dieną? • Mažiau nei 4 val. • 4-8 val.

• 9-12 val. • Daugiau nei 12 val. Kuriais socialiniais tinklais naudojats? Galite pažymėti ne vieną (<i>respondentai galėjo žymėti ne vieną atsakymą</i>) • Facebook • Instagram • WhatsApp • X • Pinterest • YouTube • LinkedIn • TikTok • Socialiniais tinklais nesinaudoju • Kita
2 klausimų blokas. Skalėje nuo 1 (visiškai nesutinku) iki 5 (visiškai sutinku) įvertinkite teiginius, atspindinčius jūsų nuomonę. <i>Skalės reikšmės: 1 – visiškai nesutinku; 2 – nesutinku; 3 – nei sutinku, nei nesutinku; 4 – sutinku; 5 – visiškai sutinku.</i> • Aš galiu lengvai atpažinti, ar socialinio tinklo turinys buvo sukurtas DI ar žmogaus • Jei žinau, kas yra turinio autorius, esu linkęs labiau pasitikėti turiniu
3 klausimų blokas. Skalėje nuo 1 (visiškai nesutinku) iki 5 (visiškai sutinku) įvertinkite teiginius, atspindinčius jūsų nuomonę. <i>Skalės reikšmės: 1 – visiškai nesutinku; 2 – nesutinku; 3 – nei sutinku, nei nesutinku; 4 – sutinku; 5 – visiškai sutinku.</i> • Žvelgiu į DI kaip į bendraautorį, padedantį perteikti informaciją • Žmogus, pasitelkdamas DI įrankius, kuria kokybiškesnius įrašus socialiniams tinklams • Manau, jog DI yra lygiavertis partneris kuriant turinį • DI yra įrankis, kuriuo naudojasi žmogus • Manau, kad žmogus turi tiesiogiai kontroliuoti, kokį turinį kuria DI
Informacinis blokas. Dabar būsite pakviesti susipažinti su tipiniu įrašu, kurį galima rasti socialiniuose tinkluose. Peržvelkite jį ir atsakykite į keletą klausimų.
1 scenarijus – įrašas sukurtas su dirbtiniu intelektu. • Vertimas į lietuvių kalbą (<i>paspaudus šį mygtuką, angliškas tekstas patampa lietuvišku</i>) EN: Just got back from the Future Tech Summit 2025 in Amsterdam, and I'm still processing what I witnessed. Seeing humanoid robots like X1 Sophia and NEO-5 in action wasn't just fascinating — it was slightly surreal. These aren't sci-fi anymore; they're shaking hands, holding conversations, and even reading emotions. One highlight was the panel with Dr. Ayanna Howard and Kohei Ogawa, where they explored what it means for robots to understand context, not just commands. The quote that stuck with me came from Ogawa-san: "A humanoid robot should not just move like us — it should understand why we move." I also had a great chat with Lucia Martens from Ethics Lab Europe about the importance of embedding empathy into robotic design, not just code. Grateful for the chance to experience this intersection of engineering and philosophy. If you're curious about where human-robot interaction is headed, this is the time to start paying attention. #robotics #futureofwork #AI

LT:

Ką tik grįžau iš „Ateities technologijų viršūnių susitikimo 2025“ Amsterdame ir vis dar apdoroju tai, ką mačiau. Stebėti humanoidinius robotus, tokius kaip X1 Sophia ir NEO-5, veikiančius gyvai, buvo ne tik žavu – bet ir šiek tiek siurrealu. Tai nebėra mokslinė fantastika; jie spaudžia rankas, palaiko pokalbius ir netgi skaito emocijas.

Vienas iš pagrindinių įvykių buvo diskusija su daktare Ayanna Howard ir Kohei Ogawa, kur jie tyrinėjo, ką reikia robotams suprasti kontekstą, o ne tik komandas. Man įstrigo Ogawa citata: „Humanoidinis robotas turėtų ne tik judėti kaip mes – jis turėtų suprasti, kodėl mes judame.“

Taip pat puikiai pasikalbėjau su Lucia Martens iš „Ethics Lab Europe“ apie empatijos įterpimo į robotų dizainą, o ne tik kodą, svarbą.

Esu dėkingas už galimybę patirti šią inžinerijos ir filosofijos sankirtą. Jei domitės, kur link krypsta žmogaus ir robotų sąveika, dabar yra laikas pradėti atkreipti dėmesį.

#robotics #futureofwork #AI



2 scenarijus – įrašas sukurtas žmogaus.

- Vertimas į lietuvių kalbą (*paspaudus šį mygtuką, angliškas tekstas patampa lietuvišku*)

EN:

Humanoids are here and they're ready to transform industries.

AI-driven humanoid robots are no longer science fiction, as I experienced firsthand at the International Humanoid Forum. Companies like Booster Robotics and Unitree Robotics demonstrated advanced humanoids that are quickly moving from research labs to real-world applications.

The market is evolving fast. While Goldman Sachs predicted a \$38B humanoid robotics market by 2035, Figure AI alone has already surpassed that valuation. As industries face skilled labour shortages, humanoids are stepping in to automate increasingly complex tasks and augment human capabilities.

With experts expecting humanoid robots to reach market maturity within the next decade, "humanoids will work alongside humans", as Xavier Comtesse emphasized. Greater Zurich is well-positioned to be a key player in this industry, with strengths in AI, traditional robotics, and a thriving ecosystem of research and industry collaboration.

#robotics #futureofwork #AI



LT:

Žmogiškieji robotai jau čia ir jie pasirengę keisti pramonės šakas.

Dirbtinio intelekto valdomi humanoidiniai robotai nebėra mokslinė fantastika, tuo įsitikinau asmeniškai Tarptautiniame humanoidų forume. Tokios kompanijos kaip „Booster Robotics“ ir „Unitree Robotics“ pademonstravo pažangius humanoidus, kurie sparčiai pereina iš tyrimų laboratorijų į realaus pasaulio pritaikymą.

Rinka sparčiai vystosi. Nors „Goldman Sachs“ prognozavo, kad iki 2035 m. humanoidinės robotikos rinka pasieks 38 mlrd. JAV dolerių, vien tik „Figure AI“ jau viršijo šią vertę. Pramonės šakoms susiduriant su kvalifikuotos darbo jėgos trūkumu, humanoidai žengia į priekį, kad automatizuotų vis sudėtingesnes užduotis ir sustiprintų žmogaus galimybes.

Ekspertams prognozuojant, kad humanoidiniai robotai rinkos brandą pasieks per ateinančią dešimtmetį, „humanoidai dirbs kartu su žmonėmis“, kaip pabrėžė Xavier Comtesse. Didysis Ciurichas yra gerai pasirengęs tapti pagrindiniu šios pramonės žaidėju, turėdamas stiprių dirbtinio intelekto, tradicinės robotikos srityse ir klestinčią tyrimų bei pramonės bendradarbiavimo ekosistemą.

#robotics #futureofwork #AI

4 klausimų blokas.

Skalėje nuo 1 (visiškai nesutinku) iki 5 (visiškai sutinku) įvertinkite teiginius, atspindinčius jūsų nuomonę apie įrašą.
Skalės reikšmės: 1 – visiškai nesutinku; 2 – nesutinku; 3 – nei sutinku, nei nesutinku; 4 – sutinku; 5 – visiškai sutinku.

- Įrašas atrodo nuoširdus
- Įrašė manipuliuojama statistiniais faktais
- Įrašė nėra faktinių klaidų
- Įrašė išvelgiu žmogiškosios patirties
- Įrašas demonstruoja gyvybingumą ir turi savą asmenybę
- Įrašas pasižymi nuoseklumu
- Įrašo formatas ir struktūra yra nuspėjami
- Įrašas yra šiurpus
- Įrašas yra neįprastas
- Įrašas, kaip vaizdo ir teksto visuma, sudaro įspūdį, jog yra sudėliotas iš nesusijusių dedamųjų
- Kai kurios įrašo detalės atrodo nenatūraliai
- Įrašas pasižymi sintetiškumu
- Įrašas, kaip visuma, sudaro nesuderintų derinių įspūdį
- Įrašas atrodo faktiškai teisingas
- Papildomai tikrinčiau pagrindinius teiginius, pateiktus įrašė
- Man buvo lengva išlaikyti dėmesį skaitant įrašą
- Įrašą peržvelgiau nuo pradžios iki pabaigos
- Įrašas yra įdomus
- Įrašas atkreipia dėmesį į kultūrinį kontekstą
- Įrašas orientuotas į temas, kurios aktualios dabar
- Įrašas turi išliekamosios vertės
- Įrašą buvo lengva suprasti

• Įrašas pateikiamas logiškai ir struktūruotai • Įrašė pastebėjau dar negirdėtų minčių • Įrašo autoriumi galima pasitikėti • Įrašas yra patikimas • Įrašo autorius yra žmogus
5.1 klausimų blokas. 1 scenarijus – įrašas sukurtas su dirbtiniu intelektu.
Šį įrašą sugeneravo dirbtinis intelektas. Kaip pasikeitė įrašo patikimumo vertinimas, sužinojus tikrąjį turinio autorių? • Įrašu pasitikiu daug labiau • Įrašu pasitikiu šiek tiek labiau • Pasitikėjimo vertinimas nepasikeitė • Įrašu pasitikiu šiek tiek mažiau • Įrašu pasitikiu daug mažiau
5.2 klausimų blokas. 2 scenarijus – įrašas sukurtas žmogaus.
Šį įrašą parašė žmogus. Kaip pasikeitė įrašo patikimumo vertinimas, sužinojus tikrąjį turinio autorių? • Įrašu pasitikiu daug labiau • Įrašu pasitikiu šiek tiek labiau • Pasitikėjimo vertinimas nepasikeitė • Įrašu pasitikiu šiek tiek mažiau • Įrašu pasitikiu daug mažiau
6 klausimų blokas.
Skalėje nuo 1 (visiškai nesutinku) iki 5 (visiškai sutinku) įvertinkite teiginius, atspindinčius jūsų ketinimus sąveikauti su turiniu. <i>Skalės reikšmės: 1 – visiškai nesutinku; 2 – nesutinku; 3 – nei sutinku, nei nesutinku; 4 – sutinku; 5 – visiškai sutinku.</i> • Tikėtina, jog perskaityčiau šį įrašą socialiniuose tinkluose • Tikėtina, jog apsilankyčiau įrašo autoriaus paskyroje • Tikėtina, jog pasidalinčiau šiuo įrašu su draugais socialiniuose • Tikėtina, jog sureaguočiau į šį turinį paspausdama (-s) „patinka“
7 klausimų blokas.
Jūsų lytis • Vyras • Moteris • Kita Jūsų amžius • Šioje vietoje prašoma respondentų amžių įvesti ranka Išsilavinimas • Pagrindinis • Vidurinis • Profesinis • Aukštasis (neuniversitetinis) • Aukštasis (universitetinis) Kaip įvertintumėte savo žinias apie didelės kalbos modelius (LLM), tokius kaip „ChatGPT“, „Gemini“, „Perplexity“, „Midjourney“ ar panašius, skalėje nuo 1 iki 5, kur 1 – visiškai neturiu žinių, 5 – ekspertinis lygis? • 1

- 2
- 3
- 4
- 5

Kaip dažnai naudojate didelių kalbos modelių (LLM) pagrindu veikiančias programėles ar sistemas, pvz. „ChatGPT“, „Gemini“, „Perplexity“, „Midjourney“ ar panašius?

- Kasdien
- Keletą kartų per savaitę
- Kartą per savaitę
- Kartą per mėnesį
- Dar rečiau
- Niekada
- Nesu susipažinęs (-usi) su šiomis programomis

Baigiamosios nuostatos.

Ačiū už Jūsų laiką!

Jeigu turite papildomų klausimų arba norite susipažinti su tyrimo rezultatais, susisiekite el. paštu airidas.kundrotas@ktu.edu

3 priedas. „ChatGPT“ ir „Gemini“ užklausa

„ChatGPT“

Užklausa

Write a LinkedIn-style post reflecting, after a technology conference in Amsterdam, on how humanoid robots are changing our understanding of humanity, consciousness, and empathy. Include a specific example from a discussion, a quote from a scientist, and mention a meeting with an ethics researcher who talks about the importance of empathy. It should sound inspiring, but realistic. It should be about 815 characters excluding spaces.

Atsakymas	Here's a LinkedIn-style post that reflects on your experience at a technology conference in Amsterdam, focusing on the deeper implications of humanoid robots: Just got back from the Future Tech Summit 2025 in Amsterdam, and I'm still processing what I witnessed. Seeing humanoid robots like X1 Sophia and NEO-5 in action wasn't just fascinating — it was slightly surreal. These aren't sci-fi anymore; they're shaking hands, holding conversations, and even reading emotions. One highlight was the panel with Dr. Ayanna Howard and Kohei Ogawa, where they explored what it means for robots to understand context, not just commands. The quote that stuck with me came from Ogawa-san: "A humanoid robot should not just move like us — it should understand why we move." <i>"I'm here. Tell me more,"</i> the room fell silent. It was uncanny—not because it was perfect, but because it was close enough to make us question the boundaries of consciousness and care. I also had a great chat with Lucia Martens from Ethics Lab Europe about the importance of embedding empathy into robotic design, not just code. Grateful for the chance to experience this intersection of engineering and philosophy. If you're curious about where human-robot interaction is headed, this is the time to start paying attention. #robotics #futureofwork #AI Let me know if you want to tailor it more to a specific audience—tech professionals, ethicists, or your own network.
„Gemini“	
Užklausa	A highly advanced humanoid robot walking confidently on a glossy exhibition floor at a high-tech robotics convention. The robot has a futuristic design with sleek white and black armor, glowing teal lights on its helmet and joints, and a mechanical yet agile structure. The background shows a crowd of people observing various technology booths and bright LED screens with tech branding. Industrial ceiling with lights and cables above. Cinematic angle, high detail, photorealistic, 4K resolution.
Atsakymas	Sure, here is an image of a humanoid robot on a glossy floor at a high-tech robotics convention: 

4 priedas. Įrašas apie anketą socialiniuose tinkluose

 Kurią nuotrauką sukūrė žmogus, o kurią – dirbtinis intelektas?

Dirbtinis intelektas jau seniai nebe mokslinė fantastika – jis rašo tekstus, kuria vizualus, generuoja vaizdo įrašus, muziką ir netgi... sąskaitas bei čekius. Kiekvieną dieną socialiniuose tinkluose ir žiniasklaidoje vis sunkiau atskirti, kas yra tikra, o kas – sukurta dirbtinio intelekto pagalba.

Ar suvokę, kad turinys sukurtas DI, mes iš karto tampame atsargesni? O gal atvirkščiai – kai turinys atrodo patikimas, mes nebekreipiame dėmesio, kas jį sukūrė?

Šiuos ir kitus klausimus nagrinėju savo baigiamojo magistro darbo tyrime. Ir kviečiu jus prisijungti!

Apklausa užtruks iki 12 minučių – pasitikrinkite, ar atskiriate DI generuotą turinį nuo žmogaus sukurto.

➡ Nuoroda į apklausą: <https://2ly.link/25s08>

P. S. tikruosius nuotraukų autorius sužinosite užpildę anketą 💡



5 priedas. Respondentų demografiniai duomenys

Scenarijus

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Pirmas scenarijus (DI)	232	50,3	50,3	50,3
	Antras scenarijus (HU)	229	49,7	49,7	100,0
	Total	461	100,0	100,0	

Jūsų lytis

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Vyras	204	44,3	44,3	44,3
	Moteris	256	55,5	55,5	99,8
	Kita	1	,2	,2	100,0
	Total	461	100,0	100,0	

Išsilavinimas

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Pagrindinis	30	6,5	6,5	6,5
	Vidurinis	47	10,2	10,2	16,7
	Profesinis	40	8,7	8,7	25,4
	Aukštasis (neuniversitetinis)	119	25,8	25,8	51,2
	Aukštasis (universitetinis)	225	48,8	48,8	100,0
	Total	461	100,0	100,0	

Kaip dažnai naudojate didelių kalbos modelių (LLM) pagrindu veikiančias programėles ar sistemas, pvz. „ChatGPT“, „Gemini“, „Perplexity“, „Midjourney“ ar panašius?

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Kasdien	70	15,2	15,2	15,2
	Keletą kartų per savaitę	118	25,6	25,6	40,8
	Kartą per savaitę	54	11,7	11,7	52,5
	Kartą per mėnesį	65	14,1	14,1	66,6
	Dar rečiau	73	15,8	15,8	82,4
	Niekada	43	9,3	9,3	91,8
	Nesu susipažinęs (-usi) su šiomis programomis	38	8,2	8,2	100,0
	Total	461	100,0	100,0	

Jūsų amžius

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	18	8	1,7	1,7	1,7
	19	12	2,6	2,6	4,3
	20	12	2,6	2,6	6,9
	21	4	,9	,9	7,8
	22	20	4,3	4,3	12,1
	23	22	4,8	4,8	16,9
	24	12	2,6	2,6	19,5
	25	19	4,1	4,1	23,6
	26	8	1,7	1,7	25,4
	27	11	2,4	2,4	27,8
	28	9	2,0	2,0	29,7
	29	16	3,5	3,5	33,2
	30	11	2,4	2,4	35,6
	31	11	2,4	2,4	38,0
	32	22	4,8	4,8	42,7
	33	12	2,6	2,6	45,3
	34	10	2,2	2,2	47,5
	35	15	3,3	3,3	50,8
	36	11	2,4	2,4	53,1
	37	15	3,3	3,3	56,4
	38	18	3,9	3,9	60,3
	39	13	2,8	2,8	63,1
	40	13	2,8	2,8	65,9
	41	10	2,2	2,2	68,1
	42	15	3,3	3,3	71,4
	43	14	3,0	3,0	74,4
	44	14	3,0	3,0	77,4
	45	17	3,7	3,7	81,1
	46	7	1,5	1,5	82,6
	47	4	,9	,9	83,5
	48	2	,4	,4	83,9
	49	4	,9	,9	84,8
	50	7	1,5	1,5	86,3
	51	8	1,7	1,7	88,1
	52	5	1,1	1,1	89,2
	53	6	1,3	1,3	90,5
	54	4	,9	,9	91,3
	55	6	1,3	1,3	92,6
	56	5	1,1	1,1	93,7
	57	3	,7	,7	94,4
	58	5	1,1	1,1	95,4
	59	3	,7	,7	96,1
	60	4	,9	,9	97,0
	61	1	,2	,2	97,2
	62	1	,2	,2	97,4
	63	3	,7	,7	98,0
	64	1	,2	,2	98,3
	65	1	,2	,2	98,5
	66	1	,2	,2	98,7
	67	4	,9	,9	99,6
	73	1	,2	,2	99,8
	74	1	,2	,2	100,0
	Total	461	100,0	100,0	

Statistics

Jūsų amžius

N	Valid	461
	Missing	0
Mean		36,25
Median		35,00
Std. Deviation		11,811
Variance		139,492
Minimum		18
Maximum		74
Percentiles	25	26,00
	50	35,00
	75	44,00

Jūsų amžius (kategorijos)

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	<= 26 m.	117	25,4	25,4	25,4
	27 – 35 m.	117	25,4	25,4	50,8
	36 – 44 m.	123	26,7	26,7	77,4
	>= 45 m.	104	22,6	22,6	100,0
	Total	461	100,0	100,0	

Kaip dažnai naudojate socialinius tinklus?

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Keletą kartų per dieną	292	63,3	63,3	63,3
	Kartą per dieną	58	12,6	12,6	75,9
	Keletą kartų per savaitę	46	10,0	10,0	85,9
	Kartą per savaitę	22	4,8	4,8	90,7
	Keletą kartų per mėnesį	43	9,3	9,3	100,0
	Total	461	100,0	100,0	

Kiek laiko praleidžiate naudodamiesi socialiniais tinklais per dieną?

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Mažiau nei 4 val.	224	48,6	48,6	48,6
	4-8 val.	141	30,6	30,6	79,2
	9-12 val.	67	14,5	14,5	93,7
	Daugiau nei 12 val.	29	6,3	6,3	100,0
	Total	461	100,0	100,0	

\$sT_Type Frequencies

		Responses		Percent of Cases
		N	Percent	
\$sT_Type ^a	Kuriais socialiniais tinklais naudojātės? Galite pažymėti ne vieną (Facebook)	300	19,1%	65,1%
	Kuriais socialiniais tinklais naudojātės? Galite pažymėti ne vieną (Instagram)	275	17,5%	59,7%
	Kuriais socialiniais tinklais naudojātės? Galite pažymėti ne vieną (WhatsApp)	149	9,5%	32,3%
	Kuriais socialiniais tinklais naudojātės? Galite pažymėti ne vieną (X)	101	6,4%	21,9%
	Kuriais socialiniais tinklais naudojātės? Galite pažymėti ne vieną (Pinterest)	170	10,8%	36,9%
	Kuriais socialiniais tinklais naudojātės? Galite pažymėti ne vieną (YouTube)	257	16,4%	55,7%
	Kuriais socialiniais tinklais naudojātės? Galite pažymėti ne vieną (LinkedIn)	151	9,6%	32,8%
	Kuriais socialiniais tinklais naudojātės? Galite pažymėti ne vieną (TikTok)	159	10,1%	34,5%
	Kuriais socialiniais tinklais naudojātės? Galite pažymėti ne vieną (Kita)	7	0,4%	1,5%
	Total	1569	100,0%	340,3%

a. Dichotomy group tabulated at value 1.

6 priedas. Tyrimo konstrukto ir skalių kokybės vertinimas

KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		,851
Bartlett's Test of Sphericity	Approx. Chi-Square	4895,016
	df	465
	Sig.	<,001

Communalities

	Initial	Extraction
Žvelgiu į DI kaip į bendraautorių, padedantį perteikti informaciją	,351	,437
Žmogus, pasitelkdamas DI įrankius, kuria kokybiškesnius įrašus socialiniams tinklams	,368	,481
Manau, jog DI yra lygiavertis partneris kuriant turinį	,345	,422
DI yra įrankis, kuriuo naudojasi žmogus	,331	,388
Manau, kad žmogus turi tiesiogiai kontroliuoti, kokį turinį kuria DI	,263	,337
Įrašas atrodo nuoširdus	,406	,373
Įrašė manipuliuojama statistiniais faktais (r)	,268	,249
Įrašė nėra faktinių klaidų	,231	,222
Įrašė įžvelgiu tikrosios žmogiškosios patirties	,427	,434
Įrašas demonstruoja gyvybingumą ir turi savą asmenybę	,591	,726
Įrašas pasižymi nuoseklumu	,523	,626
Įrašo formatas ir struktūra yra nuspėjami	,313	,311
Įrašas yra šiurpus (r)	,353	,393
Įrašas yra neįprastas (r)	,417	,444
Įrašas, kaip vaizdo ir teksto visuma, sudaro įspūdį, jog yra sudėliotas iš nesusijusių dedamųjų (r)	,496	,512
Kai kurios įrašo detalės atrodo nenatūraliai (r)	,518	,568
Įrašas pasižymi sintetiškumu (r)	,448	,626
Įrašas, kaip visuma, sudaro nesusuderintų derinių įspūdį (r)	,507	,591
Įrašas atrodo faktiškai teisingas	,241	,229
Papildomai tikrinčiau pagrindinius teiginius, pateiktus įrašė (r)	,285	,285
Man buvo lengva išlaikyti dėmesį skaitant įrašą	,559	,563
Įrašą peržvelgiau nuo pradžios iki pabaigos	,372	,467
Įrašas yra įdomus	,509	,543
Įrašas atkreipia dėmesį į kultūrinį kontekstą	,377	,316
Įrašas orientuotas į temas, kurios aktualios dabar	,470	,467
Įrašas turi išliekamosios vertės	,416	,384
Įrašė pastebėjau dar negirdėtų minčių	,337	,237
Įrašą buvo lengva suprasti	,565	,523
Įrašas pateikiamas logiškai ir struktūruotai	,521	,552
Įrašo autoriumi galima pasitikėti	,479	,535
Įrašas yra patikimas	,475	,567

Extraction Method: Principal Axis Factoring.

Total Variance Explained

Factor	Initial Eigenvalues			Extraction Sums of Squared Loadings			Rotation Sums of Squared Loadings		
	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	6,749	21,771	21,771	6,243	20,138	20,138	3,946	12,729	12,729
2	3,403	10,979	32,750	2,866	9,245	29,384	2,321	7,487	20,217
3	2,344	7,561	40,310	1,790	5,774	35,158	1,794	5,789	26,005
4	1,648	5,317	45,627	1,106	3,567	38,724	1,598	5,154	31,160
5	1,235	3,983	49,610	,713	2,301	41,025	1,525	4,919	36,078
6	1,099	3,545	53,155	,589	1,900	42,925	1,380	4,450	40,529
7	1,062	3,427	56,582	,501	1,617	44,542	1,244	4,013	44,542
8	,991	3,196	59,778						
9	,926	2,987	62,766						
10	,874	2,819	65,584						
11	,848	2,735	68,319						
12	,825	2,661	70,980						
13	,788	2,541	73,522						
14	,725	2,340	75,861						
15	,660	2,130	77,992						
16	,615	1,983	79,974						
17	,595	1,921	81,895						
18	,555	1,792	83,687						
19	,516	1,664	85,351						
20	,507	1,635	86,986						
21	,487	1,571	88,557						
22	,457	1,474	90,031						
23	,425	1,371	91,402						
24	,410	1,321	92,723						
25	,396	1,277	94,000						
26	,373	1,204	95,205						
27	,357	1,151	96,355						
28	,321	1,036	97,391						
29	,294	,949	98,340						
30	,267	,862	99,202						
31	,247	,798	100,000						

Extraction Method: Principal Axis Factoring.

Rotated Factor Matrix^a

	Factor						
	1	2	3	4	5	6	7
Žvelgiu į DI kaip į bendraautorijų, padedantį perteikti informaciją	,112	-,054	-,009	,260	,078	,577	,122
Žmogus, pasitelkdamas DI įrankius, kuria kokybiškesnius įrašus socialiniams tinklams	,267	-,026	,105	,011	-,146	,606	,097
Manau, jog DI yra lygiavertis partneris kuriant turinį	,178	,005	,080	,056	,109	,584	-,166
DI yra įrankis, kuriuo naudojasi žmogus	,056	,075	,140	,543	-,115	,157	,163
Manau, kad žmogus turi tiesiogiai kontroliuoti, kokį turinį kuria DI	,045	-,080	-,053	,554	-,125	,051	,017
Įrašas atrodo nuoširdus	,549	,109	,120	,096	,066	,159	,075
Įrašė manipuliuojama statistiniais faktais (r)	,070	,438	,001	-,086	,211	-,008	,002
Įrašė nėra faktinių klaidų	,318	,023	,230	,227	-,017	,125	-,017
Įrašė įžvelgiu tikrosios žmogiškosios patirties	,596	,136	,009	,215	,009	,110	,039
Įrašas demonstruoja gyvybingumą ir turi savą asmenybę	,794	-,067	,116	,103	,201	,141	-,088
Įrašas pasižymi nuoseklumu	,457	,228	,484	,340	-,111	,031	-,040
Įrašo formatas ir struktūra yra nuspėjami	,122	,089	,263	,305	-,351	-,062	,007
Įrašas yra šiurpus (r)	-,050	,586	-,038	,051	,028	,121	,167
Įrašas yra nejprastas (r)	-,035	,635	-,114	,089	-,073	-,117	-,003
Įrašas, kaip vaizdo ir teksto visuma, sudaro įspūdį, jog yra sudėliotas iš nesusijusių dedamųjų (r)	,159	,649	,155	-,136	,130	-,074	,001
Kai kurios įrašo detalės atrodo nenatūraliai (r)	,206	,334	,146	-,303	,530	,010	-,138
Įrašas pasižymi sintetiškumu (r)	,122	,183	-,003	-,144	,744	,049	,027
Įrašas, kaip visuma, sudaro nesuderintų derinių įspūdį (r)	,111	,575	,216	-,015	,445	,030	-,044
Įrašas atrodo faktiškai teisingas	,197	-,082	,230	,271	,023	,111	,211
Papildomai tikrinčiau pagrindinius teiginius, pateiktus įrašė (r)	-,046	,175	,123	-,354	,308	-,082	-,104
Man buvo lengva išlaikyti dėmesį skaitant įrašą	,593	,130	,261	,016	,055	-,045	,348
Įrašą peržvelgiau nuo pradžios iki pabaigos	,202	,009	,042	,298	-,082	-,004	,573
Įrašas yra įdomus	,512	,130	,256	-,077	,081	,094	,421
Įrašas atkreipia dėmesį į kultūrinį kontekstą	,515	-,044	,164	-,027	,021	,115	,087
Įrašas orientuotas į temas, kurios aktualios dabar	,406	,199	,097	,302	-,161	,221	,295
Įrašas turi išliekamosios vertės	,489	-,137	,163	-,182	,046	,210	,142
Įrašė pastebėjau dar negirdėtų minčių	,318	-,259	,134	,071	-,007	,189	,101
Įrašą buvo lengva suprasti	,411	,332	,249	,127	-,050	-,035	,403
Įrašas pateikiamas logiškai ir struktūruotai	,505	,254	,335	,097	-,143	,042	,298
Įrašo autoriumi galima pasitikėti	,328	-,015	,640	-,024	,086	,085	,050
Įrašas yra patikimas	,336	-,048	,629	-,014	,068	,120	,194

Extraction Method: Principal Axis Factoring.

Rotation Method: Varimax with Kaiser Normalization.

a. Rotation converged in 11 iterations.

KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		,695
Bartlett's Test of Sphericity	Approx. Chi-Square	464,860
	df	6
	Sig.	<,001

Communalities

	Initial	Extraction
Tikėtina, jog perskaityčiau šį įrašą socialiniuose tinkluose	,217	,224
Tikėtina, jog apsilankyčiau įrašo autoriaus paskyroje	,338	,446
Tikėtina, jog pasidalinčiau šiuo įrašu su draugais socialiniuose tinkluose	,423	,509
Tikėtina, jog sureaguočiau į šį turinį paspausdama (-s) „patinka“	,431	,590

Extraction Method: Principal Axis Factoring.

Total Variance Explained

Factor	Total	Initial Eigenvalues		Extraction Sums of Squared Loadings		
		% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	2,293	57,337	57,337	1,769	44,234	44,234
2	,799	19,967	77,304			
3	,545	13,622	90,925			
4	,363	9,075	100,000			

Extraction Method: Principal Axis Factoring.

Factor Matrix^a

	Factor 1
Tikėtina, jog perskaityčiau šį įrašą socialiniuose tinkluose	,473
Tikėtina, jog apsilankyčiau įrašo autoriaus paskyroje	,668
Tikėtina, jog pasidalinčiau šiuo įrašu su draugais socialiniuose tinkluose	,714
Tikėtina, jog sureaguočiau į šį turinį paspausdama (-s) „patinka“	,768

Extraction Method: Principal Axis Factoring.

a. 1 factors extracted. 8 iterations required.

Reliability Statistics

Cronbach's Alpha	N of Items
,869	12

Reliability Statistics

Cronbach's Alpha	N of Items
,658	3

Reliability Statistics

Cronbach's Alpha	N of Items
,742	3

Reliability Statistics

Cronbach's Alpha	N of Items
,750	2

Reliability Statistics

Cronbach's Alpha	N of Items
,666	3

Reliability Statistics

Cronbach's Alpha	N of Items
,611	2

Reliability Statistics

Cronbach's Alpha	N of Items
,747	4

7 priedas. Psichometriniai matavimai

Descriptives		Statistic	Std. Error
Kako eadna naukaqate sistemsia tiznas?	Mean	1,84	,091
	95% Confidence Interval for Mean	1,72	
	Lower Bound	1,96	
	Upper Bound		
	5% Trimmed Mean	1,71	
	Median	1,80	
	Variance	1,42	
	Std. Deviation	1,320	
	Minimum	1	
	Maximum	5	
	Range	4	
	Interquartile Range	1	
	Skewness	1,369	,114
	Kurtosis	,588	,227
Kiek laiku patalidate naudotumies sistemos dienais per diena?	Mean	1,79	,043
	95% Confidence Interval for Mean	1,70	
	Lower Bound	1,87	
	Upper Bound		
	5% Trimmed Mean	1,71	
	Median	2,00	
	Variance	,809	
	Std. Deviation	,918	
	Minimum	1	
	Maximum	4	
	Range	3	
	Interquartile Range	1	
	Skewness	,833	,114
	Kurtosis	-,110	,227
PO_Obimaisaizika, kor a	Mean	3,1519	,0082
	95% Confidence Interval for Mean	3,0815	
	Lower Bound	3,2222	
	Upper Bound		
	5% Trimmed Mean	3,1727	
	Median	3,3333	
	Variance	,581	
	Std. Deviation	,7692	
	Minimum	1,00	
	Maximum	5,00	
	Range	4,00	
	Interquartile Range	1,00	
	Skewness	-,482	,114
	Kurtosis	,159	,227
PO_Obimaisaizika, bendr	Mean	4,1584	,0034
	95% Confidence Interval for Mean	4,0880	
	Lower Bound	4,2389	
	Upper Bound		
	5% Trimmed Mean	4,3411	
	Median	4,0000	
	Variance	,809	
	Std. Deviation	,7892	
	Minimum	1,00	
	Maximum	5,00	
	Range	4,00	
	Interquartile Range	,50	
	Skewness	-,1381	,114
	Kurtosis	2,487	,227
At_V_bendr	Mean	3,1384	,0050
	95% Confidence Interval for Mean	3,0533	
	Lower Bound	3,1634	
	Upper Bound		
	5% Trimmed Mean	3,1135	
	Median	3,0833	
	Variance	,915	
	Std. Deviation	,9541	
	Minimum	1,00	
	Maximum	4,83	
	Range	3,83	
	Interquartile Range	,75	
	Skewness	-,382	,114
	Kurtosis	,581	,227
At_Risetas_bendr	Mean	2,5119	,0042
	95% Confidence Interval for Mean	2,4424	
	Lower Bound	2,5814	
	Upper Bound		
	5% Trimmed Mean	2,5318	
	Median	2,5000	
	Variance	,881	
	Std. Deviation	,9384	
	Minimum	1,00	
	Maximum	5,00	
	Range	4,00	
	Interquartile Range	1,00	
	Skewness	,293	,114
	Kurtosis	-,485	,227
Istot_bendr	Mean	2,9082	,0040
	95% Confidence Interval for Mean	2,8333	
	Lower Bound	2,9831	
	Upper Bound		
	5% Trimmed Mean	2,9181	
	Median	3,0000	
	Variance	,535	
	Std. Deviation	,7315	
	Minimum	1,00	
	Maximum	5,00	
	Range	4,00	
	Interquartile Range	,50	
	Skewness	-,346	,114
	Kurtosis	,800	,227
Istot_Veistumas_Judicinas a	Mean	2,6369	,0076
	95% Confidence Interval for Mean	2,5285	
	Lower Bound	2,7453	
	Upper Bound		
	5% Trimmed Mean	2,5938	
	Median	2,6587	
	Variance	,823	
	Std. Deviation	,7892	
	Minimum	1,00	
	Maximum	5,00	
	Range	4,00	
	Interquartile Range	1,00	
	Skewness	,114	,114
	Kurtosis	-,287	,227
Istot_Veistumas_Judicinas bendr	Mean	3,1620	,0066
	95% Confidence Interval for Mean	3,0919	
	Lower Bound	3,2330	
	Upper Bound		
	5% Trimmed Mean	3,1855	
	Median	3,3333	
	Variance	,588	
	Std. Deviation	,7692	
	Minimum	1,00	
	Maximum	4,83	
	Range	3,83	
	Interquartile Range	1,00	
	Skewness	-,488	,114
	Kurtosis	-,130	,227
Alyq_ambius	Mean	35,17	,056
	95% Confidence Interval for Mean	35,17	
	Lower Bound	37,33	
	Upper Bound		
	5% Trimmed Mean	35,72	
	Median	35,00	
	Variance	139,482	
	Std. Deviation	11,811	
	Minimum	18	
	Maximum	74	
	Range	56	
	Interquartile Range	18	
	Skewness	,588	,114
	Kurtosis	-,220	,227
Kadambiusas	Mean	4,00	,058
	95% Confidence Interval for Mean	3,89	
	Lower Bound	4,12	
	Upper Bound		
	5% Trimmed Mean	4,11	
	Median	4,00	
	Variance	1,572	
	Std. Deviation	1,254	
	Minimum	1	
	Maximum	5	
	Range	4	
	Interquartile Range	2	
	Skewness	-,114	,114
	Kurtosis	,981	,227
Kapo_pardumams_kapo Dienas apie dideles kaibus rodentes (L.M. kaibo kapo „ChatGPT“, „Gemini“, „Perplexity“, „Midjourney“ ir panalauka, skaitu nuo 1 iki 5, kur 1 – mažiausia reikšmė, 5nuo, 5 – didžiausia reikšmė?	Mean	2,60	,056
	95% Confidence Interval for Mean	2,49	
	Lower Bound	2,71	
	Upper Bound		
	5% Trimmed Mean	2,68	
	Median	3,00	
	Variance	1,414	
	Std. Deviation	1,189	
	Minimum	1	
	Maximum	5	
	Range	4	
	Interquartile Range	1	
	Skewness	,159	,114
	Kurtosis	-,981	,227
Kapo_pardumams_kapo dideles kaibus rodentes (L.M. kaibo kapo „ChatGPT“, „Gemini“, „Perplexity“, „Midjourney“ ir panalauka)	Mean	3,31	,086
	95% Confidence Interval for Mean	3,34	
	Lower Bound	3,68	
	Upper Bound		
	5% Trimmed Mean	3,65	
	Median	3,00	
	Variance	2,545	
	Std. Deviation	1,593	
	Minimum	1	
	Maximum	7	
	Range	6	
	Interquartile Range	3	
	Skewness	,333	,114
	Kurtosis	-,383	,227
Kapo_pardumams_kapo publikuoms reikšmės, skaitu nuo 1 iki 5, kur 1 – mažiausia reikšmė, 5nuo, 5 – didžiausia reikšmė?	Mean	3,13	,048
	95% Confidence Interval for Mean	3,03	
	Lower Bound	3,22	
	Upper Bound		
	5% Trimmed Mean	3,14	
	Median	3,00	
	Variance	1,381	
	Std. Deviation	1,180	
	Minimum	1	
	Maximum	5	
	Range	4	
	Interquartile Range	1	
	Skewness	-,385	,114
	Kurtosis	-,260	,227

Ranks

	Jūsų lytis	N	Mean Rank	Sum of Ranks
PO_Dlbendraautorius_bendr	Vyras	204	226,45	46196,00
	Moteris	256	233,73	59834,00
	Total	460		
PO_Dlirankis_bendr	Vyras	204	228,00	46511,50
	Moteris	256	232,49	59518,50
	Total	460		
AK_VI_bendr	Vyras	204	234,12	47760,50
	Moteris	256	227,62	58269,50
	Total	460		
AK_Elgsena_bendr	Vyras	204	208,04	42440,00
	Moteris	256	248,40	63590,00
	Total	460		
trust_bendr	Vyras	204	225,67	46036,50
	Moteris	256	234,35	59993,50
	Total	460		
benr_keistumas_justincase	Vyras	204	217,09	44286,00
	Moteris	256	241,19	61744,00
	Total	460		
benr_sintetiskumas_justincase	Vyras	204	221,82	45251,50
	Moteris	256	237,42	60778,50
	Total	460		
Kaip įvertintumėte savo žinias apie didelės kalbos modelius (LLM), tokius kaip „ChatGPT“, „Gemini“, „Perplexity“, „Midjourney“ ar panašius, skalėje nuo 1 iki 5, kur 1 – visiškai neturiu žinių, 5 – ekspertinis lygis?	Vyras	204	224,76	45851,00
	Moteris	256	235,07	60179,00
	Total	460		
Aš galiu lengvai atpažinti, ar socialinio tinklo turinys buvo sukurtas DI ar žmogaus	Vyras	204	243,97	49770,00
	Moteris	256	219,77	56260,00
	Total	460		
Jei žinau, kas yra turinio autorius, esu linkęs labiau juo pasitikėti	Vyras	204	233,17	47566,50
	Moteris	256	228,37	58463,50
	Total	460		
Kaip pasikeitė įrašo patikimumo vertinimas, sužinojus tikrąjį turinio autorių?	Vyras	204	220,10	44901,00
	Moteris	256	238,79	61129,00
	Total	460		

Test Statistics ^a											
	PO_Dlbendraa utorius_bendr	PO_Dlirankis_ bendr	AK_VI_bendr	AK_Elgsena_b endr	trust_bendr	benr_keistuma s_justincase	benr_sintetisk umas_justinca se	Kaip įvertintumėte savo žinias apie didelės kalbos modelius (LLM), tokius kaip „ChatGPT“, „Gemini“, „Perplexity“, „Midjourney“ ar panašius, skalėje nuo 1 iki 5, kur 1 – visiškai neturiu žinių, 5 – ekspertinis lygis?	Aš galiu lengvai atpažinti, ar socialinio tinklo turinys buvo sukurtas DI ar žmogaus	Jei žinau, kas yra turinio autorius, esu linkęs labiau juo pasitikėti	Kaip pasikeitė įrašo patikimumo vertinimas, sužinojus tikrąjį turinio autorį?
Mann-Whitney U	25286,000	25601,500	25373,500	21530,000	25126,500	23376,000	24341,500	24941,000	23364,000	25567,500	23991,000
Wilcoxon W	46196,000	46511,500	58269,500	42440,000	46036,500	44286,000	45251,500	45851,000	56260,000	58463,500	44901,000
Z	-,589	-,371	-,522	-,3249	-,733	-,1948	-,1262	-,852	-,2015	-,412	-,1579
Asymp. Sig. (2-tailed)	,556	,711	,602	,001	,464	,051	,207	,394	,044	,681	,114

a. Grouping Variable: Jūsų lytis

Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Kaip dažnai naudojate socialinius tinklus?	,372	461	<,001	,669	461	<,001
Kiek laiko praleidžiate naudodamiesi socialiniais tinklais per dieną?	,290	461	<,001	,782	461	<,001
PO_Dlbendraaautorius_bendr	,131	461	<,001	,970	461	<,001
PO_Dlirankis_bendr	,202	461	<,001	,849	461	<,001
AK_VI_bendr	,072	461	<,001	,988	461	<,001
AK_Elgsena_bendr	,108	461	<,001	,977	461	<,001
trust_bendr	,248	461	<,001	,916	461	<,001
benr_keistumas_justincase	,097	461	<,001	,979	461	<,001
benr_sintetiskumas_justincase	,130	461	<,001	,963	461	<,001
Jūsų amžius	,068	461	<,001	,963	461	<,001
Išsilavinimas	,275	461	<,001	,766	461	<,001
Kaip įvertintumėte savo žinias apie didelės kalbos modelius (LLM), tokius kaip „ChatGPT“, „Gemini“, „Perplexity“, „Midjourney“ ar panašius, skalėje nuo 1 iki 5, kur 1 – visiškai neturiu žinių, 5 – ekspertinis lygis?	,171	461	<,001	,900	461	<,001
Kaip dažnai naudojate didelių kalbos modelių (LLM) pagrindu veikiančias programas ar sistemas, pvz. „ChatGPT“, „Gemini“, „Perplexity“, „Midjourney“ ar panašius?	,196	461	<,001	,911	461	<,001
Kaip pasikeitė įrašo patikimumo vertinimas, sužinojus tikrąjį turinio autorių?	,219	461	<,001	,904	461	<,001

a. Lilliefors Significance Correction

Correlations																		
		Jūsų amžius	Išsilavinimas	Kaip dažnai naudojate didelių kalbos modelių (LLM) pagrindu veikiančias programas ar sistemas, pvz., ChatGPT, Gemini, Perplexity, Midjourney ar panašius?	Kaip dažnai naudojate socialinius tinklus?	Kiek laiko praleidiate naudodamiesi socialiniais tinklais per dieną?	PO_Dibendraa utoriūs_bendr	PO_Dilrankis_bendr	AK_VI_bendr	benr_keistumas JUSTINCAS	benr_sintetiskumas JUSTINCAS	trust_bendr	AK_Elgsena_bendr	Kaip įvertinūmėte savo žinias apie didelių kalbos modelių (LLM) tokius kaip ChatGPT, Gemini, Perplexity, Midjourney ar panašius? (1-5, kur 1 – visiškai neturiu žinių, 5 – ekspertinis lygis?)	Aš galiu lengvai atpažinti, ar socialinio tinklo turinys buvo sukurtas DI ar žmogaus	Jei žinau, kas yra turinio autorius, esu linkęs labiau jį patikėti		
Spearman's rho	Correlation Coefficient	1,000	,084	,160	,083	-,038	,158	,019	,047	,083	-,007	,060	,015	,161**	-,205*	-,210*	,096	
	Sig. (2-tailed)		,071	<,001	,073	,419	<,001	,681	,313	,075	,877	,201	,752	<,001	<,001	<,001	,039	
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
Išsilavinimas	Correlation Coefficient	,084	1,000	-,053	-,126*	-,087	,058	,036	-,041	-,054	,081	-,091*	,027	,017	,009	,057	,035	
	Sig. (2-tailed)		,071	,258	,007	,061	,212	,440	,384	,247	,084	,050	,557	,712	,852	,224	,451	
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
Kaip dažnai naudojate didelių kalbos modelių (LLM) pagrindu veikiančias programas ar sistemas, pvz., ChatGPT, Gemini, Perplexity, Midjourney ar panašius?	Correlation Coefficient	-,160**	-,053	1,000	-,092*	-,098*	-,079	-,086	-,097*	,101*	,024	-,050	,044	,042	-,250*	-,070	,040	
	Sig. (2-tailed)		<,001	,258		,049	,035	,091	,065	,038	,030	,613	,280	,343	,369	<,001	,135	,392
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
Kaip dažnai naudojate socialinius tinklus?	Correlation Coefficient	,083	-,126**	-,092*	1,000	-,215**	-,023	-,004	-,010	-,019	,003	,096*	-,059	-,041	,060	-,132**	-,031	
	Sig. (2-tailed)		,073	,007	,049		<,001	,615	,926	,825	,688	,956	,040	,204	,383	,202	,004	,501
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
Kiek laiko praleidiate naudodamiesi socialiniais tinklais per dieną?	Correlation Coefficient	-,038	-,087	-,098*	-,215**	1,000	,010	-,004	-,005	-,003	-,011	,080	,041	,021	,003	,002	-,016	
	Sig. (2-tailed)		,419	,061	,035	<,001	,827	,931	,916	,954	,821	,085	,384	,660	,033	,772	,724	
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
PO_Dibendraa utoriūs_bendr	Correlation Coefficient	,158**	,058	-,079	-,023	,010	1,000	,045	,303**	,026	-,044	,162**	-,015	,211**	-,010	,066	,071	
	Sig. (2-tailed)		<,001	,212	,091	,615	,827	,340	<,001	,579	,350	<,001	,750	<,001	,837	,160	,127	
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
PO_Dilrankis_bendr	Correlation Coefficient	,019	,036	-,086	-,004	-,004	,045	1,000	,107*	-,055	,101*	-,057	-,115*	-,029	,041	,111*	,258**	
	Sig. (2-tailed)		,681	,440	,065	,926	,931	,340		,022	,242	,030	,218	,014	,533	,384	<,001	
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
AK_VI_bendr	Correlation Coefficient	,047	-,041	-,097*	-,010	-,005	,303**	,107*	1,000	-,226*	-,294**	,509**	-,075	,421**	-,078	-,016	-,028	
	Sig. (2-tailed)		,313	,384	,038	,825	,916	<,001	,822		<,001	<,001	<,001	,109	<,001	,066	,571	
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
benr_keistumas JUSTINCAS	Correlation Coefficient	,083	-,054	,101*	-,019	-,003	,026	-,055	-,226**	1,000	,412**	-,036	,032	,080	-,056	-,102*	,076	
	Sig. (2-tailed)		,075	,247	,030	,688	,954	,579	,242	<,001	,443	,488	,085	,227	,029	,029	,102	
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
benr_sintetiskumas JUSTINCAS	Correlation Coefficient	-,007	,081	,024	,003	,011	-,044	,101*	,294**	,412**	1,000	-,214**	,043	-,187*	-,044	,032	,139*	
	Sig. (2-tailed)		,877	,084	,613	,956	,821	,350	,030	<,001	<,001	<,001	,356	<,001	,350	,495	,003	
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
trust_bendr	Correlation Coefficient	,060	-,091*	-,050	,096*	,080	,182**	-,057	,509**	-,036	-,214**	1,000	-,050	,246**	-,079	-,028	-,091	
	Sig. (2-tailed)		,201	,050	,280	,040	,085	<,001	,218	<,001	,443	<,001	,288	<,001	,092	,546	,050	
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
Kaip pasikeitė jūsų patikimumo vertinimas, sužinojus tikrąjį turinio autorių?	Correlation Coefficient	,015	,027	,044	-,059	,041	-,015	-,115*	-,075	,032	,043	-,050	1,000	,188*	,093	,032	-,067	
	Sig. (2-tailed)		,752	,657	,043	,204	,384	,750	,014	,109	,488	,356	,288		<,001	,075	,498	,152
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
AK_Elgsena_bendr	Correlation Coefficient	-,161**	,017	,042	-,041	,021	,215**	-,029	,421**	,080	-,183*	,246**	,168**	1,000	-,035	-,075	,055	
	Sig. (2-tailed)		<,001	,712	,369	,383	,660	<,001	,533	<,001	,085	<,001	<,001	<,001		,456	,110	,240
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	
Kaip įvertinūmėte savo žinias apie didelių kalbos modelių (LLM) tokius kaip ChatGPT, Gemini, Perplexity, Midjourney ar panašius?	Correlation Coefficient	-,205**	,009	-,250*	,060	,100*	-,010	,041	-,078	-,056	-,044	-,079	,093	-,035	1,000	,116*	-,044	
	Sig. (2-tailed)		<,001	,852	<,001	,302	,933	,837	,384	,086	,227	,350	,092	,076	,456		,313	,348
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461
Aš galiu lengvai atpažinti, ar socialinio tinklo turinys buvo sukurtas DI ar žmogaus	Correlation Coefficient	-,210**	,057	-,070	-,132**	-,014	,066	,111*	-,016	-,102*	,032	-,028	,032	-,075	,116*	1,000	-,011	
	Sig. (2-tailed)		<,001	,224	,135	,004	,772	,160	,017	,736	,029	,495	,546	,498	,110	,013		,811
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461
Jei žinau, kas yra turinio autorius, esu linkęs labiau jį patikėti	Correlation Coefficient	,096*	,036	,040	-,031	-,016	,071	,206*	-,026	,076	,198**	-,091	-,067	,095	-,044	-,045	1,000	
	Sig. (2-tailed)		,039	,451	,392	,501	,724	,127	<,001	,571	,102	,003	,050	,152	,240	,392		,811
	N	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461	461

*** Correlation is significant at the 0.01 level (2-tailed).

* Correlation is significant at the 0.05 level (2-tailed).

** Correlation is significant at the 0.01 level (2-tailed).

* Correlation is significant at the 0.05 level (2-tailed).

8 priedas. Koreliacinė kintamųjų tarpusavio ryšių analizė

Correlations													
			PO_Dibendraa utoriūs_bendr	PO_Dilrankis_ bendr	AK_VI_bendr	benr_keistuma s JUSTINCAS	benr_sintetisk umas JUSTINC AS	trust_bendr	Kaip pasikeitė jūsų patikimumo vertinimas, sužinojus tikrąjį turinio autorių?	AK_Elgsena_b_ endr	Kaip dažnai naudojate didelių kalbos modelių (LLM) pagrindu veikiančias programėles ar sistemas, pvz. „ChatGPT“, „Gemini“, „Perplexity“, „Midjourney“ ar panašius?	Aš galiu atpažinti, ar socialinio tinklo turinys buvo sukurtas DI ar žmogaus	Jei žinau, kas yra turinio autorius, esu linkęs labiau jį patikėti
Spearman's rho	PO_Dibendraa utoriūs_bendr	Correlation Coefficient	1,000	,045	,303**	,026	-,044	,182**	-,015	,211**	-,079	,066	,071
		Sig. (2-tailed)	.	,340	<.001	,579	,350	<.001	,750	<.001	,091	,160	,127
		N	461	461	461	461	461	461	461	461	461	461	461
	PO_Dilrankis_bendr	Correlation Coefficient	,045	1,000	,107*	-,055	,101*	-,057	-,115*	-,029	-,086	,111*	,256**
		Sig. (2-tailed)	,340	.	,022	,242	,030	,218	,014	,533	,065	,017	<.001
		N	461	461	461	461	461	461	461	461	461	461	461
	AK_VI_bendr	Correlation Coefficient	,303**	,107*	1,000	-,220**	-,294**	,509**	-,075	,421**	-,097*	-,016	-,026
		Sig. (2-tailed)	<.001	,022	.	<.001	<.001	<.001	,109	<.001	,038	,736	,571
		N	461	461	461	461	461	461	461	461	461	461	461
	benr_keistumas JUSTINCAS	Correlation Coefficient	,026	-,055	-,220**	1,000	,412**	-,036	,032	,080	,101*	-,102*	,076
		Sig. (2-tailed)	,579	,242	<.001	.	<.001	,443	,488	,085	,030	,029	,102
		N	461	461	461	461	461	461	461	461	461	461	461
	benr_sintetiskumas JUSTINCAS	Correlation Coefficient	-,044	,101*	-,294**	,412**	1,000	-,214**	,043	-,183*	,024	,032	,138**
		Sig. (2-tailed)	,350	,030	<.001	<.001	.	<.001	,356	<.001	,613	,495	,003
		N	461	461	461	461	461	461	461	461	461	461	461
	trust_bendr	Correlation Coefficient	,182**	-,057	,509**	-,036	-,214**	1,000	-,050	,246**	-,050	-,028	-,091
		Sig. (2-tailed)	<.001	,218	<.001	,443	<.001	.	,288	<.001	,280	,546	,050
		N	461	461	461	461	461	461	461	461	461	461	461
	Kaip pasikeitė jūsų patikimumo vertinimas, sužinojus tikrąjį turinio autorių?	Correlation Coefficient	-,015	-,115*	-,075	,032	,043	-,050	1,000	,188**	,044	,032	-,067
		Sig. (2-tailed)	,750	,014	,109	,488	,356	,288	.	<.001	,343	,498	,152
		N	461	461	461	461	461	461	461	461	461	461	461
	AK_Elgsena_bendr	Correlation Coefficient	,211**	-,029	,421**	,080	-,183*	,246**	,188**	1,000	,042	-,075	,055
		Sig. (2-tailed)	<.001	,533	<.001	,085	<.001	<.001	<.001	.	,369	,110	,240
		N	461	461	461	461	461	461	461	461	461	461	461
Kaip dažnai naudojate didelių kalbos modelių (LLM) pagrindu veikiančias programėles ar sistemas, pvz. „ChatGPT“, „Gemini“, „Perplexity“, „Midjourney“ ar panašius?	Correlation Coefficient	-,079	-,086	-,097*	,101*	,024	-,050	,044	,042	1,000	-,070	,040	
	Sig. (2-tailed)	,091	,065	,038	,030	,613	,280	,343	,369	.	,135	,392	
	N	461	461	461	461	461	461	461	461	461	461	461	
Aš galiu lengvai atpažinti, ar socialinio tinklo turinys buvo sukurtas DI ar žmogaus	Correlation Coefficient	,066	,111*	-,016	-,102*	,032	-,028	,032	-,075	-,070	1,000	-,011	
	Sig. (2-tailed)	,160	,017	,736	,029	,495	,546	,498	,110	,135	.	,811	
	N	461	461	461	461	461	461	461	461	461	461	461	
Jei žinau, kas yra turinio autorius, esu linkęs labiau jį patikėti	Correlation Coefficient	,071	,256**	-,026	,076	,138**	-,091	-,067	,055	,040	-,011	1,000	
	Sig. (2-tailed)	,127	<.001	,571	,102	,003	,050	,152	,240	,392	,811	.	
	N	461	461	461	461	461	461	461	461	461	461	461	

9 priedas. Regresinė analizė

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	,595 ^a	,354	,349	,58976

a. Predictors: (Constant), benr_sintetiskumas_justincase, AK_VI_bendr, benr_keistumas_justincase

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	86,956	3	28,985	83,334	<,001 ^b
	Residual	158,954	457	,348		
	Total	245,910	460			

a. Dependent Variable: trust_bendr

b. Predictors: (Constant), benr_sintetiskumas_justincase, AK_VI_bendr, benr_keistumas_justincase

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Collinearity Statistics	
		B	Std. Error	Beta			Tolerance	VIF
1	(Constant)	,815	,217		3,760	<,001		
	AK_VI_bendr	,680	,046	,570	14,667	<,001	,938	1,066
	benr_keistumas_justincase	,141	,038	,152	3,694	<,001	,831	1,203
	benr_sintetiskumas_justincase	-,125	,040	-,131	-3,139	,002	,810	1,234

a. Dependent Variable: trust_bendr

10 priedas. Neparametrinis Mann-Whitney U testas

Ranks

	Scenarijus	N	Mean Rank	Sum of Ranks
Kaip pasikeitė įrašo patikimumo vertinimas, sužinojus tikrąjį turinio autorių?	Pirmas scenarijus (DI)	232	171,51	39789,50
	Antras scenarijus (HU)	229	291,27	66701,50
	Total	461		

Test Statistics^a

Kaip pasikeitė įrašo patikimumo vertinimas, sužinojus tikrąjį turinio autorių?	
Mann-Whitney U	12761,500
Wilcoxon W	39789,500
Z	-10,179
Asymp. Sig. (2-tailed)	<,001

a. Grouping Variable: Scenarijus

Ranks

	Scenarijus	N	Mean Rank	Sum of Ranks
trust_bendr	Pirmas scenarijus (DI)	232	252,37	58549,00
	Antras scenarijus (HU)	229	209,35	47942,00
	Total	461		

Test Statistics^a

	trust_bendr
Mann-Whitney U	21607,000
Wilcoxon W	47942,000
Z	-3,653
Asymp. Sig. (2-tailed)	<,001

a. Grouping Variable:
Scenarijus

Ranks

	Scenarijus	N	Mean Rank	Sum of Ranks
AK_VI_bendr	Pirmas scenarijus (DI)	232	258,81	60043,50
	Antras scenarijus (HU)	229	202,83	46447,50
	Total	461		
benr_keistumas_justincas_e	Pirmas scenarijus (DI)	232	219,80	50992,50
	Antras scenarijus (HU)	229	242,35	55498,50
	Total	461		
benr_sintetiskumas_justin case	Pirmas scenarijus (DI)	232	213,62	49560,50
	Antras scenarijus (HU)	229	248,60	56930,50
	Total	461		

Test Statistics^a

	AK_VI_bendr	benr_keistumas_justincas_e	benr_sintetiskumas_justin case
Mann-Whitney U	20112,500	23964,500	22532,500
Wilcoxon W	46447,500	50992,500	49560,500
Z	-4,516	-1,833	-2,846
Asymp. Sig. (2-tailed)	<,001	,067	,004

a. Grouping Variable: Scenarijus

11 priedas. Paprastoji regresinė analizė

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	,207 ^a	,043	,041	,80794

a. Predictors: (Constant), Kaip pasikeitė krašto patikimumo vertinimas, sužinojus tikrąjį turinio autorių?

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	13,416	1	13,416	20,553	<,001 ^b
	Residual	299,624	459	,653		
	Total	313,040	460			

a. Dependent Variable: AK_Elgsena_bendr

b. Predictors: (Constant), Kaip pasikeitė įrašo patikimumo vertinimas, sužinojus tikrąjį turinio autorių?

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.	Collinearity Statistics	
		B	Std. Error	Beta			Tolerance	VIF
1	(Constant)	2,004	,119		16,786	<,001		
	Kaip pasikeitė įrašo patikimumo vertinimas, sužinojus tikrąjį turinio autorių?	,164	,036	,207	4,534	<,001	1,000	1,000

a. Dependent Variable: AK_Elgsena_bendr

Model Summary

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	,077 ^a	,006	,004	1,038

a. Predictors: (Constant), Jei žinau, kas yra turinio autorius, esu linkęs labiau juo pasitikėti

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	2,967	1	2,967	2,754	,098 ^b
	Residual	494,482	459	1,077		
	Total	497,449	460			

a. Dependent Variable: Kaip pasikeitė įrašo patikimumo vertinimas, sužinojus tikrąjį turinio autorių?

b. Predictors: (Constant), Jei žinau, kas yra turinio autorius, esu linkęs labiau juo pasitikėti

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.
		B	Std. Error	Beta		
1	(Constant)	3,456	,203		16,989	<,001
	Jei žinau, kas yra turinio autorius, esu linkęs labiau juo pasitikėti	-,081	,049	-,077	-1,660	,098

a. Dependent Variable: Kaip pasikeitė įrašo patikimumo vertinimas, sužinojus tikrąjį turinio autorių?

12 priedas. Moderacijos analizė

OUTCOME VARIABLE:

TRUST_ch

Model Summary

R	R-sq	MSE	F	df1	df2	p
,4895	,2396	,8277	47,9956	3,0000	457,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	1,4402	,7308	1,9707	,0494	,0040	2,8764
GR	1,5253	,4602	3,3143	,0010	,6209	2,4297
PO_DIira	,0515	,1725	,2988	,7652	-,2875	,3906
Int_1	-,1301	,1087	-1,1966	,2321	-,3438	,0836

OUTCOME VARIABLE:

TRUST_ch

Model Summary

R	R-sq	MSE	F	df1	df2	p
,4895	,2396	,8277	47,9956	3,0000	457,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	1,4402	,7308	1,9707	,0494	,0040	2,8764
GR	1,5253	,4602	3,3143	,0010	,6209	2,4297
PO_DIira	,0515	,1725	,2988	,7652	-,2875	,3906
Int_1	-,1301	,1087	-1,1966	,2321	-,3438	,0836

13 priedas. Chi-kvadrato testas

Teisingaidentifikuota * Žinios apie DI įrankius (kategorijos) Crosstabulation

Count

		Žinios apie DI įrankius (kategorijos)			Total
		Mažos	Vidutinės	Didelės	
Teisingaidentifikuota	Teisingai identifikuota	69	46	35	150
	Neteisingai identifikuota	143	91	77	311
Total		212	137	112	461

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)
Pearson Chi-Square	,152 ^a	2	,927
Likelihood Ratio	,152	2	,927
Linear-by-Linear Association	,032	1	,858
N of Valid Cases	461		

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 36,44.

Symmetric Measures

		Value	Approximate Significance
Nominal by Nominal	Phi	,018	,927
	Cramer's V	,018	,927
N of Valid Cases		461	

Irašo autorius yra žmogus

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	visiškai nesutinku	43	9,3	9,3	9,3
	nesutinku	126	27,3	27,3	36,7
	nei sutinku, nei nesutinku	183	39,7	39,7	76,4
	sutinku	99	21,5	21,5	97,8
	visiškai sutinku	10	2,2	2,2	100,0
	Total	461	100,0	100,0	

Teisingai identifikuota * Manau, jog galiu lengvai atskirti DI turinį nuo žmogaus sukurto Crosstabulation

Count

		Manau, jog galiu lengvai atskirti DI turinį nuo žmogaus sukurto			Total
		mano, kad negali tinkamai identifikuoti turinio autoriaus	nežino ar gali tinkamai identifikuoti turinio autorių	mano, kad gali tinkamai identifikuoti turinio autorių	
Teisingai identifikuota	Teisingai identifikuota	39	41	70	150
	Neteisingai identifikuota	75	110	126	311
Total		114	151	196	461

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)
Pearson Chi-Square	3,041 ^a	2	,219
Likelihood Ratio	3,091	2	,213
Linear-by-Linear Association	,287	1	,592
N of Valid Cases	461		

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 37,09.

Symmetric Measures

		Value	Approximate Significance
Nominal by Nominal	Phi	,081	,219
	Cramer's V	,081	,219
N of Valid Cases		461	