



Kauno technologijos universitetas

Ekonomikos ir verslo fakultetas

Dirbtinio intelekto sukurto rinkodaros turinio ir autorystės žymėjimo poveikis vartotojų atsakui socialiniuose tinkluose

Baigiamasis magistro projektas

Kristina Grybaitė

Projekto autorė

Doc. dr. Beata Šeinauskienė

Vadovė

Kaunas, 2025



Kauno technologijos universitetas

Ekonomikos ir verslo fakultetas

Dirbtinio intelekto sukurto rinkodaros turinio ir autorystės žymėjimo poveikis vartotojų atsakui socialiniuose tinkluose

Baigiamasis magistro projektas
Marketingo valdymas (kodas 6211LX038)

Kristina Grybaitė

Projekto autorė

Doc. dr. Beata

Šeinauskienė

Vadovė

Prof. dr. Žaneta

Gravelines

Recenzentė

Kaunas, 2025



Kauno technologijos universitetas

Ekonomikos ir verslo fakultetas

Kristina Grybaitė

Dirbtinio intelekto sukurto rinkodaros turinio ir autorystės žymėjimo poveikis vartotojų atsakui socialiniuose tinkluose

Akademinių 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.

Kristina Grybaitė

Patvirtinta elektroniniu būdu

Grybaitė, Kristina. Dirbtinio intelekto sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikis vartotojų atsakui socialiniuose tinkluose. Baigiamasis magistro projektas.

Vadovė doc. dr. Beata Šeinauskienė; Kauno technologijos universitetas, Ekonomikos ir verslo fakultetas.

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

Reikšminiai žodžiai: dirbtinis intelektas, socialiniai tinklai, rinkodaros turinys, dirbtinio intelekto etiketės.

Kaunas, 2025. 114 p.

Santrauka

Pastaraisiais metais vis daugiau dėmesio skiriama dirbtinio intelekto sukurtam rinkodaros turiniui ir jo poveikiui vartotojų elgsenai skaitmeninėje aplinkoje. Tačiau mokslinėje literatūroje vis dar nesutariama, kaip dirbtinio intelekto autorystės žymėjimas veikia vartotojų atsaką. Kai kurie tyrimai rodo, kad dirbtinio intelekto sukurtas turinys gali būti vertinamas taip pat palankiai kaip ir žmogaus kurtas (Kirkby et al., 2023). Tačiau kiti autoriai pažymi, kad dirbtinio intelekto autorystės žyma dažnai sukelia skepticizmą, mažina suvokiamą autentiškumą ir neigiamai veikia turinio vertinimą (Zhang & Gosline, 2023; Lim & Schmalzle, 2024; Shantatula et al., 2024).

Šio tyrimo objektas – vartotojų atsakas į dirbtinio intelekto sukurtą rinkodaros turinį ir jo autorystės žymėjimą socialiniuose tinkluose.

Tyrimo tikslas – teoriškai ir empiriškai pagrįsti dirbtinio intelekto sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikį vartotojų atsakui socialiniuose tinkluose.

Atlikto tyrimo rezultatai patvirtina, kad dirbtinio intelekto turinio poveikis vartotojams priklauso nuo kontekstinio suderinamumo – pažymėtas dirbtinio intelekto turinys sukelia teigiamą atsaką tik tuomet, kai jo tematika yra susijusi su dirbtiniu intelektu ir kai šaltinis ir turinys vartotojui atrodo konceptualiai nuoseklūs (Lee et al., 2025). Šis efektas aiškinamas remiantis schemos atitikimo teorija, teigiančia, kad vartotojai pozityviau vertina informaciją, kuri dera su jų išankstinėmis žiniomis ar lūkesčiais (Mandler, 1982; Meyers-Levy & Tybout, 1989). Savo ruožtu, dirbtinio intelekto ir jo sukurto turinio atitikimas prognozuoja aukštesnius patikimumo vertinimus – tiek šaltinio empatijos, palankumo, kompetencijos, tiek turinio patikimumo dimensijose.

Visgi, tyrimas atskleidė, kad toks šaltinio ir turinio atitikimas neturi reikšmingo praktinio poveikio vartotojų įsitraukimo ketinimams; tik suvokiamas turinio patikimumas bei šaltinio emocinės savybės – ypač empatija ir palankumas – skatina elgseninį atsaką. Be to, semantinė analizė parodė, kad vartotojai kelia aukštesnius kokybinius ir etinius reikalavimus dirbtinio intelekto turiniui, ypač kai jis publikuojamas akademinų institucijų vardu – jie tikisi kalbinio tikslumo, stilistinio natūralumo bei skaidraus tokio turinio žymėjimo. Tyrimo įžvalgos gali prisidėti prie atsakingesnio dirbtinio intelekto sprendimų taikymo rinkodaroje ir informuoti organizacijas, siekiančias užtikrinti etišką bei pasitikėjamą stiprinančią komunikaciją socialiniuose tinkluose.

Grybaitė Kristina. The Impact of Labelling AI-Generated Marketing Content and Its Authorship on Consumer Response in Social Media. Master's Final Degree.

Supervisor Assoc. Prof. Dr Beata Šeinauskienė; School of Economy and Business, Kaunas University of Technology.

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

Keywords: artificial intelligence, social media, marketing content, artificial intelligence labels.

Kaunas, 2025. 114 pages.

Summary

In recent years, increasing attention has been devoted to marketing content generated by artificial intelligence (AI) and its impact on consumer behaviour in the digital environment. However, academic literature still lacks consensus on how the labelling of AI authorship influences consumer responses. Some studies suggest that AI-generated content may be evaluated as favourably as that created by humans (Kirkby et al., 2023). However, other authors point out that labelling content as AI-generated often triggers scepticism, reduces perceived authenticity, and negatively affects content evaluation (Zhang & Gosline, 2023; Lim & Schmalzle, 2024; Shantatula et al., 2024).

The focus of this study is on consumer response to the labelling of AI-generated marketing content and its authorship in social media.

The aim of the study is to provide a theoretical and empirical justification of the impact of labelling AI-generated marketing content and authorship on consumer response on social media.

The results of the conducted study confirm that the impact of AI-generated content on consumers depends on contextual congruence – content labelled as AI-generated elicits a positive response only when its topic relates to artificial intelligence and when the source and the content appear conceptually consistent to the consumer (Lee et al., 2025). This effect is explained by schema congruity theory, which posits that individuals respond more positively to information that aligns with their prior knowledge or expectations (Mandler, 1982; Meyers-Levy & Tybout, 1989). Consequently, the alignment between artificial intelligence and the content it generates predicts higher credibility ratings – across dimensions such as source empathy, goodwill, competence, and the perceived credibility of the content itself.

Nevertheless, the study revealed that such source–content congruence does not have a significant practical effect on consumers' engagement intentions; only perceived content credibility and the emotional characteristics of the source – particularly empathy and goodwill – drive behavioural responses. Furthermore, semantic analysis indicated that consumers place higher qualitative and ethical demands on AI-generated content, especially when it is published under the name of academic institutions – they expect linguistic accuracy, stylistic naturalness, and transparent labelling. These insights may contribute to the more responsible use of AI solutions in marketing and inform organisations seeking to ensure ethical and trust-enhancing communication on social media.

Turinys

Lentelių sąrašas	7
Paveikslų sąrašas	9
Įvadas.....	10
1. Dirbtinio intelekto sukurtos rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose tyrimų aktualumas ir problematika.....	13
2. Dirbtinio intelekto sukurtos rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose teorinis pagrindimas	23
2.1. Dirbtinio intelekto sukurtos rinkodaros turinio socialiniuose tinkluose ypatybės	23
2.1.1. Dirbtinio intelekto sukurtas tekstinis rinkodaros turinys socialiniuose tinkluose	24
2.1.2. Dirbtinio intelekto sukurtas vizualinis rinkodaros turinys socialiniuose tinkluose.....	26
2.2. Vartotojų atsaką į dirbtinio intelekto sukurtą turinio autorystę paaiškinančios teorijos	29
2.3. Vartotojų atsakas į dirbtinio intelekto sukurtą rinkodaros turinį ir jo autorystės žymėjimą	32
2.3.1. Suvokiamas šaltinio ir turinio patikimumas	36
2.3.2. Vartotojų suvokiamas šaltinio ir turinio atitikimas	39
2.3.3. Vartotojų įsitraukimas socialiniuose tinkluose.....	41
2.4. Dirbtinio intelekto sukurtos rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose konceptualus modelis.....	45
3. Dirbtinio intelekto sukurtos rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose tyrimo metodologija.....	49
3.1. Tyrimo objektas, tikslas ir uždaviniai	50
3.2. Empirinio tyrimo metodai	50
3.3. Empirinio tyrimo eigos ir duomenų analizės procedūros.....	55
4. Dirbtinio intelekto sukurtos rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose, empirinio tyrimo rezultatai.....	60
4.1. Tyrimo imties charakterizavimas pagal sociodemografinės charakteristikos	60
4.2. Tyrimo instrumento metodologinės kokybės analizė	63
4.3. Eksperimento grupių homogeniškumo ir manipuliacijų veiksmingumo analizė	71
4.4. Atvirų atsakymų, nurodančių vartotojų emocinį atsaką į manipuliacinį turinį, sentimentų analizė.....	75
4.5. Hipotezių tikrinimo rezultatai.....	78
4.6. Dirbtinio intelekto sukurtos rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose tyrimo rezultatų diskusija	96
Išvados ir rekomendacijos	101
Literatūros sąrašas	105
Informacijos šaltinių sąrašas	114
Priedai.....	115
1 priedas. ChatGPT sugeneruotas tekstas.....	115
2 priedas. Išsami mokslinių šaltinių, aprašančių vartotojų požiūrį į DI valdomų sprendimų sukurtą rinkodaros turinį, apžvalga.	116
3 priedas. Naudotų matavimo skalių ir jų adaptacijų palyginimas tyrime	127
4 priedas. Tyrimo anketa	131
5 priedas. Užklauso, pateiktos scenarijų turiniui sukurti	139
6 priedas. Tyrimo respondentų charakteristikos.....	143

7	priedas. Faktorinės analizės rezultatai. Suvokiamas šaltinio patikimumas.....	150
8	priedas. Faktorinės analizės rezultatai. Vartotojų įsitraukimo ketinimai	153
9	priedas. Faktorinės analizės rezultatai. Suvokiamas šaltinio ir turinio atitikimas.....	155
10	priedas. Faktorinės analizės rezultatai. Suvokiamas turinio patikimumas	159
11	priedas. Skalių patikimumo vertinimas. Šaltinio patikimumo per pasitikėjimą skalė	161
12	priedas. Skalių patikimumo vertinimas. Turinio patikimumo skalė.....	162
13	priedas. Skalių patikimumo vertinimas. Šaltinio patikimumo per kompetenciją skalė	163
14	priedas. Skalių patikimumo vertinimas. Šaltinio patikimumo per palankumą skalė	165
15	priedas. Skalių patikimumo vertinimas. Šaltinio patikimumo per empatiją skalė	166
16	priedas. Skalių patikimumo vertinimas. Vartotojų įsitraukimo ketinimų skalė	167
17	priedas. Skalių patikimumo vertinimas. Šaltinio ir turinio atitikimo skalė.....	168
18	priedas. Tyrimo kintamųjų aprašomoji analizė	169
19	priedas. Homogeniškumo testavimas	175
20	priedas. Manipuliacijos analizė	180
21	priedas. Atvirų atsakymų analizė	183
22	priedas. Hipotezių testavimas	192

Lentelių sąrašas

1 lentelė. Vartotojų požiūrio į dirbtinio intelekto valdomų sprendimų sukurtą rinkodaros turinį mokslinių šaltinių apžvalga.	18
2 lentelė. Vartotojų atsakas į DI sukurtą turinį ir jo autorystės atskleidimo poveikis vartotojų sprendimams.....	32
3 lentelė. Vartotojų įsitraukimo socialiniuose tinkluose sampratų palyginimas.....	42
4 lentelė. 2x2 tarpgrupinio eksperimento dizaino scenarijų struktūra.	51
5 lentelė. Tyrimo respondentų demografinės charakteristikos (N = 372).....	60
6 lentelė. Respondentų pasiskirstymas pagal socialinių tinklų naudojimo trukmę per dieną (N = 372)	61
7 lentelė. Respondentų dažniausiai naudojami socialiniai tinklai (N = 372).....	62
8 lentelė. Skalių tinkamumo vertinimas tiriamajai faktorių analizei (N = 372).....	63
9 lentelė. Šaltinio patikimumo konstrukto teiginiai pagal išskirtus faktorius (N = 372).....	64
10 lentelė. Vartotojų įsitraukimo ketinimų konstrukto teiginiai pagal išskirtus faktorius (N = 372).....	65
11 lentelė. Suvokiamo šaltinio ir turinio atitikimo konstrukto teiginiai pagal išskirtus faktorius (N = 372).....	65
12 lentelė. Suvokiamo turinio patikimumo konstrukto teiginiai pagal išskirtus faktorius (N = 372).....	66
13 lentelė. Tiriamų konstrukto skalių patikimumo vertinimas (N = 372)	66
14 lentelė. Tyrimo kintamųjų vertinimo charakteristikos (N = 372)	68
15 lentelė. Tyrimo kintamųjų pasiskirstymo analizė naudojantis Kolmogorovo–Smirnov ir Shapiro–Wilko testais (N = 372)	68
16 lentelė. Kintamųjų koreliacijos analizė (N = 372)	69
17 lentelė. Pearsono chi kvadrato testo rezultatai analizuojant respondentų pasiskirstymo eksperimentinėse grupėse homogeniškumą (N = 370).....	72
18 lentelė. Lygių dispersijų prielaidos Lavene'o testas: amžius (N = 372).....	72
19 lentelė. Išsilavinimo rangų vidurkiai pagal eksperimentines grupes	73
20 lentelė. Pearsono chi kvadrato testo rezultatai analizuojant turinio sąsajos su DI ir respondentų autorystės vertinimą (N = 372).....	73
21 lentelė. Kryžminės dažnių analizės rezultatai vertinant turinio sąsajos su DI ir respondentų autorystės sąveiką (N = 372)	74
22 lentelė. Respondentų vertinimų pasiskirstymas dėl DI įrašo autorystės pagal eksperimentines grupes (N = 372).....	74
23 lentelė. Atvirų respondentų atsakymų autorystės priskyrimo ir emocinio tono analizė (N = 180)	77
24 lentelė. Suvokiamo šaltinio ir turinio atitikimo vidurkiai pagal DI autorystės žymėjimą (N = 372)	78
25 lentelė. Lygių dispersijų prielaidos Lavene'o testas: DI autorystės žymėjimo ir turinio sąsajos su DI poveikis suvokiamam šaltinio ir turinio atitikimui (N = 372).....	79
26 lentelė. Dvifaktoriinės dispersijos analizės rezultatai: DI žymėjimo ir turinio sąsajos su DI poveikis suvokiamam šaltinio ir turinio atitikimui (N = 372)	79
27 lentelė. Suvokiamo šaltinio ir turinio atitikimo vidurkiai pagal DI žymėjimą ir turinio sąsają su DI (N = 372)	80
28 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis šaltinio kompetencijai (N = 372)	82

29 lentelė. Regresijos koeficientų analizės rezultatai, tiriant šaltinio ir turinio atitikimo poveikį šaltinio kompetencijai (N = 372).....	82
30 lentelė. Modelio suvestinė: šaltinio ir turinio atitikimo poveikis šaltinio palankumui (N = 372).....	82
31 lentelė. Regresijos koeficientų analizės rezultatai: šaltinio ir turinio atitikimo poveikis šaltinio palankumui (N = 372)	83
32 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis šaltinio empatijai (N = 372)...	83
33 lentelė. Regresijos koeficientų analizės rezultatai, tiriant šaltinio ir turinio atitikimo poveikį šaltinio empatijai (N = 372).....	83
34 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis šaltinio pasitikėjimui (N = 372)	83
35 lentelė. Regresijos koeficientų analizė: šaltinio ir turinio atitikimo poveikis šaltinio pasitikėjimui (N = 372)	84
36 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis turinio patikimumui (N = 372)	84
37 lentelė. Regresijos koeficientų analizės rezultatai: šaltinio ir turinio atitikimo poveikis turinio patikimumui (N = 372).....	85
38 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis vartotojų įsitraukimo ketinimams (N = 372)	85
39 lentelė. Regresijos koeficientų analizė: šaltinio ir turinio atitikimo poveikis įsitraukimo ketinimams (N = 372).....	86
40 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis vartotojų įsitraukimo ketinimams (N = 372)	86
41 lentelė. Pagrindiniai rezultatai: koeficientų analizė (N = 372)	87
42 lentelė. Regresijos modelių, tiriant ryšį tarp šaltinio ir turinio atitikimo, turinio patikimumo ir įsitraukimo ketinimų, rezultatai	88
43 lentelė. Tiesioginė, netiesioginė (per turinio patikimumą) ir suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams (N=372)	88
44 lentelė. Regresijos modelių, tiriant ryšį tarp šaltinio ir turinio atitikimo, šaltinio kompetencijos ir įsitraukimo ketinimų, rezultatai	89
45 lentelė. Tiesioginė, netiesioginė (per šaltinio kompetenciją) ir suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams (N=372)	89
46 lentelė. Regresijos modelių, tiriant ryšį tarp šaltinio ir turinio atitikimo, šaltinio palankumo ir įsitraukimo ketinimų, rezultatai	90
47 lentelė. Tiesioginė, netiesioginė (per šaltinio palankumą) ir suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams (N=372)	91
48 lentelė. Regresijos modelių, tiriant ryšį tarp šaltinio ir turinio atitikimo, šaltinio empatijos ir įsitraukimo ketinimų, rezultatai	92
49 lentelė. Tiesioginė, netiesioginė (per šaltinio empatiją) ir suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams (N=372)	92
50 lentelė. Regresijos modelių, tiriant ryšį tarp šaltinio ir turinio atitikimo, pasitikėjimo šaltiniu ir įsitraukimo ketinimų, rezultatai	93
51 lentelė. Tiesioginė, netiesioginė (per šaltinio empatiją) ir suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams (N=372)	93
52 lentelė. Hipotezių tikrinimo rezultatų suvestinė (N = 372).....	95

Paveikslų sąrašas

1 pav. Žymos, kuriomis žymimas DI sukurtas turinys socialinėse platformose „Facebook“, „Instagram“ ir „TikTok“	14
2 pav. Konceptualus dirbtinio intelekto sukurtos rinkodaros turinio ir autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose modelis	48
3 pav. Tyrimo instrumento schema	52
4 pav. Semantinės respondentų atvirų atsakymų analizės žodžių debesis (N = 180)	76
5 pav. Šaltinio ir turinio atitikimo vertinimų palyginimai nepriklausomų kintamųjų grupėse (N = 372)	81
6 pav. Šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams per turinio patikimumą	89
7 pav. Šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams per šaltinio kompetenciją	90
8 pav. Šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams per šaltinio palankumą	91
9 pav. Šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams per šaltinio empatiją	93
10 pav. Šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams per pasitikėjimą šaltiniu	94

Įvadas

Dirbtinis intelektas sparčiai populiarėja įvairiose srityse, kartu ir rinkodaroje, atverdamas naujas galimybes kurti personalizuotus ir efektyvius rinkodaros sprendimus: net 88 proc. rinkodaros specialistų naudoja dirbtinį intelektą atlikdami kasdienes užduotis (Survey monkey Inc., 2024). Dirbtinio intelekto technologijos plačiai naudojamos siekiant automatizuoti anksčiau didesnio darbuotojų įsitraukimo, finansinių ir laiko sąnaudų reikalavusius procesus, įskaitant tačiau neapsiribojant turinio kūrimu, rinkodaros duomenų rinkimu ir analize, rinkodaros strategijų formavimu, klientų aptarnavimu ir kt. (Kumar et al., 2019, Lee & Cho, 2020; Huang & Rust, 2021). Ypač socialinių tinklų rinkodaroje, dirbtinio intelekto technologijos padeda srities specialistams kurti patrauklų turinį, atsižvelgiant į vartotojų elgseną ir turinio vartojimo įpročius (Brüns & Meißner, 2024).

Vienas iš pagrindinių dirbtinio intelekto taikymo socialiniuose tinkluose aspektų yra turinio kūrimas naudojant didžiuosius kalbos modelius (angl. *large language models, LLM*). Vienas iš tokių technologijų pavyzdžių – „ChatGPT“, surenkantis daugiau nei 200 milijonų aktyvių naudotojų kiekvieną savaitę (Fried, 2024). Pateikus užklausas (žr. 1 priedas) „Trumpai (100 ž) paaiškink man, kaip veikia „ChatGPT?“, „Kokiu būdu, pasitelkiant tavo pagalbą, rinkodaros specialistai gali kurti rinkodaros turinį?“, „ChatGPT“ sugeneravo atsakymą, jog ši dirbtinio intelekto sistema geba analizuoti vartotojo pateikiamus tekstinius duomenis (užklausas) bei greitai parengti atsakymus, kurti skirtingus, personalizuotus tekstus ar kitokį kūrybinį turinį remiantis „tikimybiniais ryšiais tarp žodžių ir frazių“. Toks įrankio funkcionalumas yra reikšmingas socialinių tinklų rinkodaros srityje, nes jo naudojimas nereikalauja specifinių techninių įgūdžių ar programavimo žinių, įrankis yra prieinamas plačiam vartotojų ratui bei suteikia galimybę sparčiai ir dideliais kiekiais generuoti tikslinėms auditorijoms pritaikytą turinį (Pavlik, 2023; Abi-Rafeh et al., 2024). Atsižvelgiant į sparčią didžiųjų kalbos modelių integraciją į rinkodaros turinio kūrimą, kyla praktinis poreikis suprasti, kaip vartotojai suvokia dirbtinio intelekto sukurtus pranešimus socialiniuose tinkluose.

Tyrimo problema. Anksčiau rinkodaros turinys, sukurtas naudojant dirbtiniu intelektu grįstus sprendimus bei publikuojamas ar kitaip įveiklinamas viešoje erdvėje, nebuvo išskiriamas ar kitaip pažymimas. Tačiau pastaruoju metu, atsižvelgiant į teisės aktų pokyčius ir skaidraus dirbtinio intelekto naudojimo poreikį, tam tikrose socialinių tinklų platformose buvo realizuotos etiketės (angl. *labels*), nurodančios, kad įrašas buvo sukurtas pasitelkiant dirbtinio intelekto technologijas, taip didinant platformos vartotojų sąmoningumą apie matomą turinį. Panašios, rekomenduojamos vartoti etiketės jau yra taikomos socialinių tinklų platformų įrašuose, kurių publikavimas yra finansuojamas prekių ženklų ar organizacijų, pažymint įrašus kaip „remiama“ (angl. *sponsored*). Tokios etiketės, leidžiančios vartotojams atpažinti žinutes kaip reklamą, skatina nepasitikėjimą, mažina vartotojų norą dalytis pranešimu ar kitaip išreikšti pritarimą matomam turiniui (Boerman et al., 2017). Kitos įspėjamosios etiketės, kaip „klaidinga“ (angl. *false*), „pakeista“ (angl. *altered*) arba „iš dalies klaidinga“ (angl. *partily false*), „Facebook“ tinkle taikomos turiniui, kuris yra peržiūrėtas nepriklausomų faktų tikrintojų bei nustatytas kaip neatitinkantis platformos standartų, taip pat daro įtaką vartotojų atsakui. Anot tyrimų (Mena, 2020; Clayton et al., 2020; Koch et al., 2023), įspėjamosios etiketės reikšmingai sumažina vartotojų ketinimus dalytis publikuojama informacija.

Iki šiol dirbtinio intelekto sugeneruotą turinį žyminčios etiketės buvo tik rekomendacinio pobūdžio, tačiau tai nesuteikė pakankamai skaidrumo bei neužtikrino tinkamo platformų vartotojų informavimo apie turinio autorystę. Tyrimai rodo, kad dalis turinio kūrėjų neįžvelgia būtinybės deklaruoti dirbtinio

intelekto naudojimo siekiant kurti turinį, laikydami šią technologiją lygiaverte kitoms įprastoms priemonėms, tokioms kaip teksto redagavimo ar automatinio taisymo įrankiai (Draxler et al., 2024). Neužtikrinant skaidrumo, vartotojų požiūris į skaitmeninės rinkodaros kanalus, kuriuose naudojamas dirbtinis intelektas, ateityje gali reikšmingai prastėti (Campbell et al., 2022). Dėl to, kai kurios socialinių tinklų platformos 2024 m. ėmė naudoti algoritmus, galinčius automatiškai atpažinti dirbtinio intelekto sukurtą turinį ir jį pažymėti, įskaitant „Meta Platforms Inc.“ priklausančius „Facebook“, „Instagram“ ir „Threads“ (Clegg, 2024) bei „ByteDance“ priklausančią „TikTok“ (TikTok, n.d.). Šis pokytis kelia svarbius klausimus apie tai, kaip vartotojai reaguoja į dirbtinio intelekto sukurtą bei pažymėtą turinį ir ar turinio žymėjimas gali turėti įtakos jų atsakui – požiūriui į turinį naudojantį prekių ženklą, pasitikėjimui skelbiama informacija ir įsitraukimui.

2024 m. gegužę atliktas pasaulinis tyrimas parodė, kad net 62 procentai rinkodaros, viešųjų ryšių, pardavimų ir klientų aptarnavimo specialistų mano, jog dirbtinio intelekto sukurtam turiniui žymėti skirtos etiketės galimai darytų teigiamą įtaką jų socialinių tinklų strategijai (Dixon, 2024a). Tačiau tų pačių metų gegužę Jungtinėje Karalystėje atlikta apklausa parodė, kad net 27 proc. respondentų, pamatę socialiniuose tinkluose pažymėtą dirbtinio intelekto sukurtą turinį, blokuotų arba nebesektų paskyros, kuri jį paskelbė (Dixon, 2024b). Išsiskiriančias vartotojų reakcijas patvirtina ir moksliniai tyrimai. Kai kurie autoriai pabrėžia, kad vartotojai, sužinoję apie turinio sukūrimą naudojant dirbtinį intelektą, jaučia didesnę skeptiškumą ir mažesnę pasitikėjimą tokiu turiniu (Kirkby et al., 2023; Altay & Gilardi, 2024). Kiti tyrimai – priešingai, rodo, kad vartotojų požiūris į dirbtinio intelekto sugeneruotą turinį nesiskiria nuo jų požiūrio į žmogaus sukurtą turinį (Shantatula et al., 2024). Nepaisant augančio mokslinio susidomėjimo šia tema, didžioji dalis atliktų tyrimų naudojo nenatūralias ar konceptualias žymėjimo formas – autorystės informacija buvo pateikiama atskirai nuo paties turinio (Altay & Gilardi, 2024; Lee et al., 2025) ar įterpiama kaip papildomas tekstinis paaiškinimas (Carney et al., 2024). Todėl empiriniai tyrimai, analizuojantys DI sukurtų įrašų žymėjimo poveikį natūralioje, platformos struktūroje integruotoje formoje, vis dar yra itin riboti.

Be to, iki šiol mažai dėmesio skirta dar vienam svarbiam aspektui – turinio teminei sąsajai su dirbtiniu intelektu. Šis veiksnys gali turėti daryti įtaką vartotojų atsakui, nes DI žymėjimas gali būti vertinamas skirtingai priklausomai nuo to, ar pats įrašas yra apie DI grįstą produktą. Todėl didėjant dirbtinio intelekto technologijų integracijai į įvairias rinkodaros strategijas, pastebimas nuolatinis poreikis suprasti, kaip šios technologijos gali būti panaudojamos rinkodaros turiniui kurti (Taylor, 2019), koks yra vartotojų požiūris į jų naudojimą rinkodaroje (Guha et al., 2020) ir suprasti, kaip dirbtinio intelekto sukurtas turinio žymėjimas gali paveikti reklamos efektyvumą (Peres et al., 2023). Tai ypač aktualu rinkodaros specialistams, kurie turi prisitaikyti prie naujų teisės aktų ir vartotojų lūkesčių, kad galėtų kurti veiksmingas ir etiškas rinkodaros strategijas socialiniuose tinkluose.

Tyrimo problema – kaip dirbtinio intelekto sukurtas rinkodaros turinys ir jo autorystės žymėjimas veikia vartotojų atsaką socialiniuose tinkluose?

Tyrimo objektas – vartotojų atsakas į dirbtinio intelekto sukurtą rinkodaros turinį ir jo autorystės žymėjimą socialiniuose tinkluose.

Tyrimo tikslas – teoriškai ir empiriškai pagrįsti dirbtinio intelekto sukurtas rinkodaros turinio ir jo autorystės žymėjimo poveikį vartotojų atsakui socialiniuose tinkluose.

Tyrimo uždaviniai:

1. atskleisti dirbtinio intelekto sukurto rinkodaros turinio ir jo autorystės žymėjimo socialiniuose tinkluose poveikio vartotojų atsakui tyrimų aktualumą ir problematiką;
2. remiantis esamų teorijų panaudojimo galimybėmis, argumentuoti konceptualų modelį, hipotetizuojantį dirbtinio intelekto sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikį vartotojų atsakui socialiniuose tinkluose,
3. pagrįsti dirbtinio intelekto sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikį vartotojų atsakui socialiniuose tinkluose, tyrimo metodologiją,
4. empiriškai patikrinti dirbtinio intelekto sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikį vartotojų atsakui socialiniuose tinkluose.

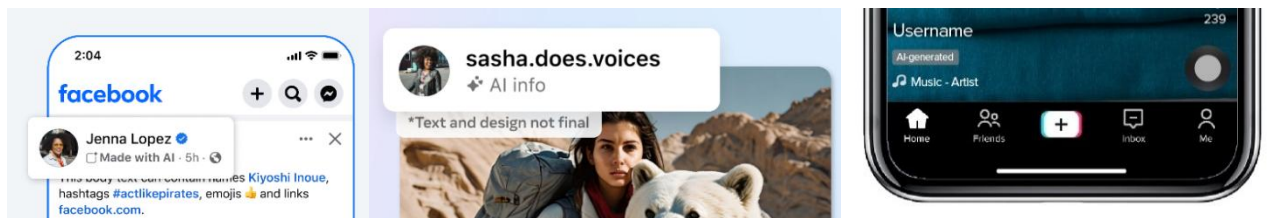
Tyrimo metodai. Tyrime taikyta sisteminė bei palyginamoji mokslinės literatūros analizė, atliktas kiekybinio pobūdžio tyrimas. Duomenys buvo renkami pasitelkiant internetinę anketinę apklausą. Statistinei duomenų analizei naudoti įvairūs metodai: aprašomoji statistika, faktorinė analizė, koreliaciniai ir regresiniai skaičiavimai, taip pat taikyti neparametriniai testai. Surinkti duomenys apdoroti naudojant statistinę analizės programinę įrangą „SPSS“.

1. Dirbtinio intelekto sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose tyrimų aktualumas ir problematika

Dirbtinio intelekto naudojimas rinkodaroje. Dirbtinis intelektas (toliau – DI) vis dažniau naudojamas pasikartojančių rinkodaros užduočių automatizavimui, duomenų apdorojimui ir analizei (Huang & Rust, 2020). Šios kompiuterinės „sistemos sugebėjimas teisingai interpretuoti išorinius duomenis, mokytis iš jų ir panaudoti tokias žinias, kad būtų įgyvendinti konkretūs tikslai ir uždaviniai, lanksčiai juos pritaikant“ (Kaplan & Haenlein, 2019, p. 15) kai yra „veikiama nepakankamų žinių ir ribotų išteklių sąlygomis“ (Wang, 2019), leidžia kurti turinį (Gao et al., 2023, Abi-Rafteh et al., 2024), analizuoti rinkodaros struktūruotus ir nestrukūruotus duomenis, įskaitant skaitinius duomenis, vaizdus ar kalbą, kurie atskleidžia gilesnius vartotojų elgsenos modelius (Davenport et al., 2019). Tokiais duomenimis grįsti rinkodaros sprendimai gali padėti optimizuoti reklamos išlaidas, efektyviai pasiekti tikslinę auditoriją (Lee & Cho, 2020) bei kurti jai pritaikytą patirtį, pavyzdžiui, pasitelkiant personalizuotus rekomendacijų algoritmus (Kumar et al., 2019).

Didėjant DI naudojimui rinkodaros srityje, stebimas ir DI grįstų sprendimų prieinamumo augimas, ypač tų, kuriems nereikalingos specializuotos technologinės žinios. Tai sudaro sąlygas rinkodaros specialistams savarankiškai naudoti pažangius įrankius, tokius kaip dideliais kalbos modeliais grįsti generatyvaus DI sprendimai, pavyzdžiui, 2022 m. „OpenAI“ sukurtas pokalbių robotas „ChatGPT“, siekiant kurti turinį įvairiems rinkodaros komunikacijos kanalams: nuo socialinių tinklų įrašų iki tinklaraščių straipsnių (Wiredu et al., 2023). „ChatGPT“ naudotojams suteikia galimybę tiesiogiai įvesti tekstines užklausas ir per itin trumpą laiką gauti išsamius, kontekstualius atsakymus (Pavlik, 2023). Jei tinkamai instrukuotas, šis generatyvaus DI modelis gali generuoti įtaigias, dideliu efektyvumu pasižyminčias personalizuotas žinutes įvairiems rinkodaros kanalams (Matz et al., 2024), įskaitant socialinius tinklus.

Anksčiau DI sukurto turinio žymėjimas socialiniuose tinkluose nebuvo privalomas, todėl vartotojai dažnai nesuvokdavo, kad sąveikauja su DI sukurtais žinutėmis. Tai lėmė tokio turinio naujumas bei jo kūrėjų požiūris į DI technologijas kaip į įprastus teksto redagavimo ar automatinio taisymo įrankius, kurių naudojimo nurodyti nėra būtinybės (Draxler et al., 2024). Tačiau skaidraus DI naudojimo poreikį socialiniuose tinkluose paskatino Europos Sąjungos DI aktas, įsigaliojęs 2024 m. rugpjūčio 1 d. Nors dokumente nėra apibrėžtas DI sukurto turinio naudojimas socialiniuose tinkluose, aktas numato, jog asmenys turi būti aiškiai informuoti, jei jie sąveikauja su DI sistema (2024 m. birželio 13 d. Europos Parlamento & Tarybos reglamentas (ES) 2024/1689, 2024). Taip siekiama užtikrinti skaidrumą, jog vartotojai suprastų, kai susiduria su DI technologijomis. Todėl socialinėse platformose imta diegti algoritmus, gebančius atpažinti DI sukurtą turinį ir jį pažymėti atitinkamomis etiketėmis (žr. 1 pav.), pavyzdžiui, „sukurta DI“ (angl. *made with AI / AI info / AI-generated*) (Clegg, 2024; TikTok, n.d.). Visgi, DI valdomų sprendimų sukurto rinkodaros turinio žymėjimas socialiniuose tinkluose kelia klausimus apie vartotojų požiūrį į tokį turinį ir jo įtaką vartotojų atsakui.



1 pav. Žymos, kuriomis žymimas DI sukurtas turinys socialinėse platformose „Facebook“, „Instagram“ ir „TikTok“

Siekiant išsamiai ištirti dirbtinio intelekto sukurtą rinkodaros turinio autorystės poveikį vartotojų atsakui socialiniuose tinkluose bei nustatyti egzistuojančių tyrimų ribotumus ir tolimesnių tyrimų perspektyvas, svarbu atlikti sistemingą esamų mokslinių tyrimų apžvalgą. Analizės pagalba nustatyti tyrimų iššūkiai leistų identifikuoti būsimų mokslinių tyrimų prioritetus. Atsižvelgiant į DI taikymo rinkodaroje mokslinių tyrimų novatoriškumą, identifikuotos dvi pagrindinės tyrimų sritys: vartotojų atsakas į turinio žymėjimą socialiniuose tinkluose bei į DI valdomų sprendimų sukurtą rinkodaros turinį.

Vartotojų atsako į turinio žymėjimą socialiniuose tinkluose ištirtumo lygmuo. Socialinių tinklų vartotojų reakcijos į turinį žyminčias etiketes dažniausiai priklauso nuo žymėjimo tipo, jo pateikimo būdo ir vartotojų suvokimo apie tokių etikečių reikšmę. Tyrimai rodo, kad turinio žymos, tokios kaip „remiama“, padeda vartotojams geriau suprasti, kad turinys su kuriuo sąveikaujama yra komercinio pobūdžio, o tai dažnai skatina kritišką požiūrį į tokį turinį (Boerman et al., 2017; De Veirman & Hudders, 2020). Pavyzdžiui, vartotojai yra linkę skeptiškiau vertinti nuomonės formuotojų ar kitų žinomų asmenų paskelbtus įrašus, kai juose yra aiškiai nurodytas įrašo rėmimo aspektas, mat tokios žymos atskleidžia turinį skelbiančių asmenų finansines paskatas bei suponuoja turinio subjektyvumą ar nepatikimumą. Anot Boerman'as et al. (2017), tokio tipo socialinių tinklų turinio žymos taip pat dažnai mažina vartotojų norą įsitraukti į turinį, pavyzdžiui, juo dalytis ar komentuoti, nes vartotojai gali jausti, kad jų veiksmai tokiu būdu prisideda prie reklaminio tikslo įgyvendinimo.

Mena'o (2020) bei Clayton'o et al. (2020) tyrimai parodė, kad įspėjamosios etiketės, tokios kaip „klaidinga“ arba „užginčytas turinys“ (angl. *disputed content*), mažina pasitikėjimą žymimu turiniu ir vartotojų tikimybę juo dalytis. Tokios žymos socialiniuose tinkluose buvo realizuotos siekiant kovoti su skleidžiama dezinformacija; jos padeda vartotojams identifikuoti galimai klaidingą informaciją ir paskatina juos kritiškiau įvertinti matomą turinį. Visgi, tyrimai rodo, jog kitų tinklo vartotojų išreikštas pritarimas turiniui gali nusverti neigiamą žymėjimo poveikį (Luo et al., 2022; Molina et al., 2023). Luo et al. (2022) tyrimo rezultatai atskleidė, kad antraštės, sulaukusios daug „Facebook“ paspaudimų „patinka“ (angl. *like*), buvo laikomos patikimesnėmis nei tos, kurios sulaukė mažiau vartotojų įsitraukimo, o Molina et al. (2023) nustatė, kad vartotojai yra labiau linkę palikti komentarus po klaidingas naujienas pristatančiais įrašais, kai jie turi daugiau „patinka“ paspaudimų. Tai rodo, jos socialiniuose tinkluose „patinka“ paspaudimų skaičius stiprina vartotojų suvokiamą turinio patikimumą.

Tačiau egzistuoja ir tam prieštaraujantys tyrimai. Remiantis Koch'u et al. (2023) tyrimo duomenimis, kai vartotojai sąveikauja su jau pažymėtu turiniu socialiniuose tinkluose, turinio teisingumą tikrinančios etiketės veiksmingai mažina klaidingo turinio pasidalijimą ir jų efektyvumas nėra veikiamas kitų vartotojų įsitraukimo – pasidalinimų ar „patinka“ paspaudimų. Todėl vienareikšmio

atsakymo, ar vartotojų įsitraukimą rodantys elementai socialiniuose tinkluose veikia kitų tinklo vartotojų požiūri į matomą turinį, nėra.

Anot tyrimų, vartotojų reakcijas taip pat gali veikti etikečių aiškumas ir informatyvumas. Clayton'as et al. (2020) atkreipė dėmesį, kad konkrečios etiketės, kaip „klaidinga“, yra veiksmingesnės nei bendro pobūdžio įspėjimai, nes jos vartotojams suteikia aiškia ir pagrįstą priežastį abejoti matomu turiniu. Kita vertus, vartotojams sunkiai suprantamos žymos gali sukelti ne tik skeptiškumą dėl žymimo turinio, bet ir bendro pasitikėjimo socialinių tinklų platformos turiniu sumažėjimą.

Apibendrinant, ankstesni tyrimai apie turinio žymėjimą socialiniuose tinkluose atskleidžia, kad įvairūs žymėjimai, tokie kaip „remiama“, „klaidinga“ ar „užginčytas turinys“, dažnai sulaukia neigiamų vartotojų reakcijų. Tokios etiketės sukelia skepticizmą, mažina pasitikėjimą pažymėtu turiniu ir skatina vartotojus vengti įsitraukimo, nes jos atskleidžia turinio komercinį ar klaidinantį pobūdį. Šie rezultatai leidžia manyti, kad panašų poveikį gali turėti ir DI naudojimą turinio kūrimo atskleidžiančios etiketės, pavyzdžiui, „sukurta DI“. Nors tokios žymos skirtos užtikrinti skaidrumą dėl turinio kilmės (Clegg, 2024; TikTok, n.d.), jos taip pat gali sukelti abejones dėl turinio autentiškumo ir tikslumo, ypač jei vartotojai neturi pakankamai žinių apie DI naudojimą (Fisher et al., 2024). Todėl nėra aišku, ar DI sukurto rinkodaros turinio autorystės žymėjimas darytų tokią pat įtaką, kaip ir kitos socialiniuose tinkluose realizuotos etiketės, ar jo poveikis vartotojų pasitikėjimui turiniu bei jų įsitraukimui socialiniuose tinkluose reikšmingai skirtųsi. Atsižvelgiant į tai, kad turinio žymėjimas nebėra savanoriškas, svarbu tęsti mokslinius tyrimus, siekiant išsiaiškinti vartotojų atsaką bei nustatyti, kaip toks turinio žymėjimas gali paveikti suvokiamą turinio patikimumą, tiek vartotojų ketinimus įsitraukti.

Naujausi tyrimai, nagrinėjantys DI sukurto turinio žymėjimą socialiniuose tinkluose, nurodo, jog DI sukurtų antraščių žymėjimas tokiomis etiketėmis, kaip „sukurta DI“, socialiniuose tinkluose sukelia vartotojų skepticizmą, kuris mažina tiek suvokiamą naujienos antraštės tikslumą, tiek norą ja dalytis (Altay & Gilardi, 2024). Šis poveikis pasireiškė nepriklausomai nuo to, ar antraštė buvo teisinga, klaidinga, žmogaus ar DI sukurta. Tai rodo, kad turinio žymėjimas kaip „sukurta DI“ savaime silpnina vartotojų pasitikėjimą turiniu. Šis skepticizmas gali kilti iš vartotojų klaidingo suvokimo, jog antraščių, pažymėtų kaip sukurtų DI, kūrimo procesas buvo visiškai automatizuotas ir nebuvo prižiūrimas žmogaus. Nors ši prielaida nėra teisinga, ji reikšmingai paveikia respondentų reakcijas. Altay'aus ir Gilardi'o (2024) tyrimo dalyviai dažnai laikė, kad DI sukurtas turinys yra mažiau patikimas nei žmogaus, net jei objektyviai vertinant turinio kokybę ir tikslumą buvo vienodi. Tyrimas taip pat atskleidė, kad „sukurta DI“ etiketės naudojimo poveikis buvo pastebimai mažesnis nei etiketės „klaidinga“ žymėjimo poveikis. Autoriai pažymi, jog „klaidinga“ žyma tris kartus stipriau paveikė vartotojų suvokimą ir ketinimus dalytis turiniu. Šis rezultatas rodo, kad vartotojai yra jautresni tiesioginiams teiginiams apie turinio klaidingumą nei informaciniams žymėjimams apie turinio kilmę. Tačiau DI žymėjimas vis tiek reikšmingai sumažino pasitikėjimą turiniu.

Autoriai pabrėžia, kad šie rezultatai atskleidžia gilesnį problemos aspektą – vartotojų ribotą suvokimą apie DI vaidmenį turinio kūrimo. Daugelis respondentų linkę manyti, kad DI sugeneruotas turinys yra mažiau kokybiškas ir autentiškas, nepaisant objektyvių turinio savybių. Tai, anot autorių, kelia klausimus apie DI žymėjimo politikos efektyvumą: nors skaidrumas dėl turinio kilmės yra svarbus, jis gali skatinti nepasitikėjimą, jei visuomenėje trūksta žinių apie DI veikimo principus ir kontrolės mechanizmus.

Visgi, Altay'aus ir Gilardi'o (2024) tyrime pagrindinis dėmesys buvo sutelktas į naujienų antraščių žymėjimą kaip DI sukurtų ir vartotojų reakcijas į tai – toks ribotas turinio tipas kaip naujienų antraštės nesuteikia pakankamai duomenų apie vartotojų požiūrį į tradicinę tekstinę ir vizualinę rinkodaros komunikaciją socialiniuose tinkluose. Atsižvelgiant į skirtumus tarp naujienų antraščių ir rinkodaros žinučių socialiniuose tinkluose, reikalingi papildomi tyrimai, atskleidžiantys vartotojų atsaką ir įsitraukimą socialiniuose tinkluose į DI žymėjimus, įtraukiant tiek tekstinį, tiek vizualinį turinį.

Vartotojų požiūrio į DI valdomų sprendimų sukurtą rinkodaros turinį ištirtumo lygmuo. Atliekant jau egzistuojančių mokslinių šaltinių, nagrinėjančių vartotojų požiūrį į DI sprendimų sukurtą turinį, analizę pastebima, jog šioje srityje atliktų tyrimų rezultatai neretai skiriasi. Toliau pateikiamas glaustas apibendrinimas, kuriame analizuojami skirtingų mokslinių darbų duomenys ir išvados.

Wu'o et al. (2021) straipsnyje buvo nagrinėjama, kaip žmonės 2018–2019 m. socialinėje platformoje „Twitter“ diskutuoja apie DI naudojimą reklamoje. Rezultatai parodė, jog daugiausia diskusijų sulaukė DI naudojimas socialinių tinklų kampanijose, reklamos tikslinis orientavimas bei DI valdomi įrankiai naudojami rinkodaroje. Pozityviausiai buvo vertinami DI įrankiai, skirti reklamos efektyvumui gerinti (skirti duomenų analizei, tiksliniam reklamos orientavimui ar personalizuotų reklaminių tekstų kūrimui), o neigiamai – DI naudojimas socialinių medijų kampanijose, pavyzdžiui, pakeičiant nuomonės formuotojus ar atliekant įvairius su vartotojų privatumo pažeidimais siejamus veiksmus. Autoriai pabrėžė, kad šiame tyrime išryškėjo vieningo požiūrio trūkumas apie DI naudojimą rinkodaroje.

Wu'o ir Wen'o (2021) tyrimas nagrinėjo, kaip DI naudojimas reklamų kūrime veikia vartotojų reakcijas bei požiūrį. Respondentai tyrime vertino DI sukurtas reklamas pagal jų objektyvumą, žmogiškumo lygį, keistumą (angl. *eeriness*) bei bendrą patrauklumą. Rezultatai parodė, kad DI sukurtų reklamų objektyvumas teigiamai veikė vartotojų požiūrį, nes toks turinys buvo suvokiamas kaip tikslesnis ir mažiau šališkas nei žmogaus sukurtos reklamos. Tačiau didesnis DI žmogiškumo lygis ir dėl to respondentų jaučiamas diskomfortas lėmė neigiamą poveikį jų požiūriui, sustiprinant suvokiamą reklamos keistumą ir mažinant jos patrauklumą.

Argan'as, Dinç'as et al. (2022) nagrinėjo, kaip socialinių tinklų vartotojai reaguoja į DI kuriamas reklamas. Rezultatai atskleidė tris pagrindinius atsako proceso etapus: „priėmimą“, „įsitraukimą“ ir „nutraukimo tašką“ (angl. *break-point*). Pirmame – priėmimo – etape reklama buvo vertinama kaip patraukli arba nepatraukli. Patrauklumas priklausė nuo tokio turinio aktualumo ir vizualinės kokybės. Reklamos, kurios atitiko vartotojų asmeninius interesus ar sukėlė smalsumą, buvo vertinamos teigiamai. Pavyzdžiui, reklamos, skatinančios praleidimo baimę (angl. *fear of missing-out*), efektyviai įtraukė vartotojus, nes pabrėžė ribotų pasiūlymų ar galimybių praleidimo riziką. Antrasis – įsitraukimo – etapas apėmė vartotojų gilesnę reklamos analizę. Dalyviai dažnai lygino reklamose pateiktas galimybes ir ieškojo papildomos informacijos apie paslaugas. Taip pat jie pabrėžė reklamos efektyvumą atsižvelgiant į laiko taupymo aspektą – greita ir tiksli informacija skatino vartotojus gilintis toliau. Be to, vartotojai teigiamai vertino reklamas, kurios padėjo suprasti apie kitų produktų poreikį ir juos įsigyti. Trečiasis etapas – nutraukimo taškas – įvyko, kai vartotojai nusprendė, ar tęsti sąveiką su reklama. Reklamos buvo atmetamos, jei sukeldavo nusivylimą dėl per didelio kiekio pasirinkimo galimybių, neatitikdavo vartotojų lūkesčių arba jei pateikta informacija buvo neaiški. Kai kurie dalyviai taip pat paminėjo, kad jų sprendimą veikė ankstesnė neigiama patirtis su panašiomis prekėmis ar paslaugomis.

Kirkby'is et al. (2023) tyrė, kaip DI sukurto turinio autorystės atskleidimas veikia prekių ženklo autentiškumą ir vartotojų požiūrį į prekių ženklą. Tyrime buvo nagrinėtos trys emocinio turinio kategorijos: mažai emocinis (funkcinės produkto specifikacijos), vidutiniškai emocinis (produkto aprašymas) ir aukštai emocinis (pokalbių robotų tekstas). Rezultatai parodė, kad DI sukurtos turinio autorystės atskleidimas nedarė reikšmingo poveikio prekių ženklo balso autentiškumo suvokimui ar vartotojų požiūriui į prekių ženklą. Tekstai, pažymėti kaip sukurti DI, buvo vertinami beveik identiška, kaip žmogaus sukurti tekstai. Be to, tyrimas atskleidė, kad prekių ženklo balso autentiškumas teigiamai veikia bendrą prekių ženklo autentiškumą, o šis savo ruožtu daro teigiamą įtaką vartotojų požiūriui į prekių ženklą. Nepaisant autorių iškeltų teorinių prielaidų, emocingesni tekstai, tokie kaip pokalbių robotų susirašinėjimai, taip pat neigiamai nepaveikė prekių ženklo autentiškumo suvokimo.

Visgi, egzistuoja ir prieštarų vertinimų. Zhang'o ir Gosline'ės (2023) straipsnyje buvo nagrinėjama, kaip vartotojai vertina DI, žmogaus sukurtą ir jų bendradarbiavimu paremtą turinį, bei analizuojama, ar šie vertinimai kyla dėl žmogaus favoritizmo ar dėl tiesioginio skeptiškumo DI technologijų atžvilgiu. Eksperimentuose buvo tiriamos įvairios sąlygos: kai turinio autorystė nebuvo nurodyta, kai dalyviams buvo pateikiama dalinė informacija apie autorystę ir kai autorystė buvo aiškiai pažymėta kaip žmogaus arba DI. Rezultatai atskleidė, kad žmonės nuosekliai teikė pirmenybę žmogaus kurtam turiniui, ypač kai buvo pažymėta, jog jį sukūrė žmogus, net jei DI turinys buvo objektyviai geresnis. Nepažymėtas DI turinys gavo geresnius įvertinimus nei tas, kuris buvo aiškiai pažymėtas kaip DI sukurtas. Mokslininkai tyrime daro išvadą, kad vartotojų pirmenybė žmogaus kurtam turiniui kyla iš giliai įsišaknijusių socialinių normų, kurios sieja žmogiškumą su autentiškumu ir emociniu jautrumu.

Shantatula et al. (2024) nagrinėjo, kaip DI sukurtas turinys ir vartotojų sukurtas turinys veikia socialinių tinklų rinkodaros strategijas. Tyrimo rezultatai parodė, kad tiek DI sukurtas turinys, tiek vartotojų sukurtas turinys vartotojus yra paskatinęs įsigyti atitinkamą paslaugą ar produktą. Visgi, nors DI naudojimas turiniui sukurti pasižymi galimybe greitai generuoti didelius personalizuotos medžiagos kiekius, o toks turinys respondentų dažnai buvo vertinamas kaip vidutiniškai autentiškas. Vartotojų sukurtas turinys buvo laikomas patikimesniu ir labiau susijusiu su vartotojų vertybėmis.

Lim'o ir Schmalzle'ės (2024) tyrimas, nagrinėjantis, kaip informacijos šaltinio atskleidimas veikia DI ir žmogaus sukurto turinio vertinimą sveikatos komunikacijoje, parodė, kad DI kurtas turinys buvo vertinamas geriau nei žmogaus, kai šaltinis nebuvo atskleistas, tačiau atskleidžiant autorystę reikšmingo skirtumo tarp DI ir žmogaus sukurto turinio vertinimo nebuvo – nepažymėtas DI turinys buvo vertinamas kaip patikimesnis ir efektyvesnis nei pažymėtas, tačiau ne mažiau patikimas ar efektyvus nei žmogaus. Taip pat buvo nustatyta, kad žmonės, nusistatę prieš DI, geriau vertino DI kurtą turinį, kai autorystė buvo nurodyta. Šis paradoksalus rezultatas gali būti paaiškinamas tuo, kad aiškus autorystės atskleidimas sumažina suvokiamą turinio netikrumą ir padidina pasitikėjimą, net jei respondentai yra skeptiški šaltinio atžvilgiu.

Gu'o et al. (2024) tyrime buvo nagrinėta, kaip DI sukurtų reklamų charakteristikos veikia vartotojų suvokimą ir reklamų priėmimą. Tyrimas analizavo keturis pagrindinius DI reklamų elementus: tikroviškumą, gyvybingumą, kūrybiškumą ir sintetiškumą. Rezultatai parodė, kad reklamos tikroviškumas ir kūrybiškumas mažino suvokiamą keistumą, kuris paprastai siejamas su neigiamomis emocijomis ir mažesniu pasitikėjimu. Tuo tarpu gyvybingumas ir kūrybiškumas didino suvokiamą DI intelektualumą, kuris buvo tiesiogiai susijęs su vartotojų teigiamu požiūriu į reklamas.

Sintetiškumo jausmas turėjo priešingą poveikį: jis didino keistumo pojūtį ir mažino suvokiamą intelektualumą. Tyrimo rezultatai taip pat atskleidė, kad keistumo pojūtis neigiamai veikė vartotojų norą priimti DI sukurtas reklamas, o suvoktas intelektualumas turėjo teigiamą poveikį.

Park‘as et al. (2024) tyrime buvo nagrinėjamas „Instagram“ vartotojų požiūris į skirtingų paskyrų tipus, įskaitant DI, nuomonės formuotojų ir paprastų vartotojų paskyras, bei jose skelbiamą turinį. Rezultatai parodė, kad dauguma respondentų nesugebėjo atskirti DI turinio nuo žmogaus sukurto teksto, nuotraukų ir vaizdo įrašų. Tačiau naudotojų vertinimai turinio kokybei skyrėsi priklausomai nuo paskyros tipo: nuomonės formuotojų paskyros buvo laikomos profesionaliausiomis ir patikimiausiomis, tačiau mažiau natūraliomis, o paprastų vartotojų paskyros buvo vertinamos kaip autentiškiausios, nors ir mažiau profesionalios. DI paskyros buvo laikomos inovatyviomis ir įdomiomis, panašiai kaip ir nuomonės formuotojų paskyros, tačiau jomis pasitikėta šiek tiek mažiau. Nepaisant to, vertinant paskyrų patrauklumą ir vartotojų teikiamą pirmenybę, DI paskyros buvo vertinamos taip pat teigiamai kaip ir nuomonės formuotojų paskyros, todėl autoriai išskyrė, jog ateityje DI paskyros galėtų pakeisti nuomonės formuotojų turinį.

Lee‘o et al. (2025) tyrime buvo siekiama išanalizuoti, kaip vartotojai reaguoja į DI autorystę sporto reklamose. Eksperimento dalyviams buvo rodomos DI sugeneruotos reklamos, su realiais žmonėmis, virtualiais pseudoportretais ir skaitmeniniais dvyniais – itin tikroviškomis vizualinėmis reprezentacijomis, sukurtomis pagal tikrus asmenis. Be to, buvo manipuluojama informacijos apie DI kilmę pateikimu – dalis dalyvių sužinojo apie DI autorystę prieš reklamą, dalis po reklamos, o likusiems dalyviams ši informacija nebuvo atskleista. Rezultatai parodė, kad informacijos apie DI kilmę pateikimo laikas reikšmingai paveikė vartotojų požiūrį į reklamą. Kai informacija buvo atskleista prieš peržiūrą, reklamos vertinimai buvo labiau neigiami, tačiau kai ši informacija buvo pateikta po reklamos, vartotojai reklamą vertino pozityviau.

Vartotojų požiūrio į dirbtinio intelekto valdomų sprendimų sukurtą rinkodaros turinį mokslinių šaltinių apžvalga pateikiama žemiau (žr. 1 lentelė).

1 lentelė. Vartotojų požiūrio į dirbtinio intelekto valdomų sprendimų sukurtą rinkodaros turinį mokslinių šaltinių apžvalga.

Autorius (-iai), metai	Tyrimo objektas	Tyrimo tikslai / hipotezės	Tyrimo rezultatai
Wu et al., 2021	Įrašai „Twitter“ platformoje apie DI naudojimą rinkodaroje	Išsiaiškinti, kokiomis temomis diskutuojama apie DI reklamą „Twitter“ bei kokios yra šių diskusijų nuotaikos	Pagrindinės temos: DI reklamos tikslinis orientavimas, žmogaus ir DI sąveika, DI įrankiai, DI turinio kūrimas. DI įrankiai, skirti reklamos efektyvumui didinti, buvo vertinami teigiamai. Tačiau DI naudojimas socialinių tinklų kampanijų kūrimo procese buvo vertinamas neigiamai dėl galimų privatumo pažeidimų, manipuliacijos ir autentiškumo stokos.
Wu & Wen, 2021	Veiksniai, lemiantys vartotojų požiūrį į DI sukurtas reklamas	Nustatyti, kaip reklamos kūrimo objektyvumas, DI žmogiškumo suvokimas ir vartotojų jaučiamas nejaukumas robotų atžvilgiu daro įtaką vartotojų požiūriui į DI sukurtą rinkodaros turinį. Hipotezės:	Tyrimo rezultatai parodė, kad tais atvejais, kai vartotojai mano, kad reklama sukurta objektyviai, jie ją vertina palankiau. Tačiau jei žmonės jaučia nejaukumą dėl robotų ar DI technologijų, tai gali turėti dvilypį

		(H ₁) „Mašininė euristika“ (DI kaip patikimo ir objektyvaus suvokimas) teigiamai veikia vartotojų požiūrį. (H ₂) Suvokiamas DI reklamos „keistumas“ (vartotojų diskomfortas dėl nenatūralumo) neigiamai veikia jų požiūrį. (H ₃) Suvokiamas reklamos kūrimo objektyvumas teigiamai veikia euristiką, bet mažina keistumą.	poveikį – iš vienos pusės, sustiprinti pasitikėjimą DI (kaip tikslu ir patikimu įrankiu), bet tuo pačiu ir sukelti nerimą, kuris mažina reklamos patrauklumą.
Argan et al., 2022	Vartotojų elgesys reaguojant į DI sukurtas reklamas socialinėje medijoje.	Nustatyti, kaip socialinės medijos vartotojai reaguoja į DI sukurtas reklamas. Hipotezės: (H ₁) Vartotojai teigiamai reaguoja į DI reklamas, kurios sukuria susidomėjimą, pavyzdžiui, pasitelkiant elementus, skatinančius baimę praleisti (angl. <i>fear of missing out, FoMO</i>), asmeninius interesus ar informacijos gavimą; (H ₂) DI reklamos, kurios taupo laiką ir pateikia naujas idėjas sulaukia daugiau vartotojų dėmesio; (H ₃) Reklamos, kurios neatitinka vartotojų lūkesčių arba kelia neaiškumų, yra vertinamos neigiamai.	Tyrimas atskleidė tris vartotojų reakcijų į DI reklamas etapus: priėmimas (teigiamos ir neigiamos reakcijos), ištraukimas (galimybė palyginti produktus / paslaugas neieškant papildomos informacijos, laiko taupymas, papildomų žinių suteikimas) ir galutinė nuomonė (pozityvi ar negatyvi). Reklamos, kurios atitinka vartotojų lūkesčius ir suteikia naujos, vartotojams nežinomos, bet naudingos informacijos, yra geriau vertinamos, o neatitinkančios lūkesčių ar per daug abstrakčios – atmetamos.
Kirkby et al., 2023	Prekių ženklo „Adidas“ rinkodaros tekstai, kurie atskleidžiami kaip sukurti DI, sukurti žmogaus arba jų autorystė nėra atskleidžiama.	Įvertinti, ar DI sukurtas turinio autorystės atskleidimas daro poveikį suvokiamam prekių ženklo balso autentiškumui, bendram turinio autentiškumui ir vartotojų požiūriui į jį. Hipotezės: (H _{1a}) DI autorystės atskleidimas mažina prekių ženklo balso autentiškumą; (H _{1b}) DI autorystės atskleidimas neturi įtakos prekių ženklo balso autentiškumui; (H _{2a}) DI autorystės atskleidimas mažina teigiamą požiūrį į prekių ženklą; (H _{2b}) DI šaltinio atskleidimas neturi įtakos požiūriui į prekių ženklą.	Tyrimas parodė, kad DI sukurtas turinys nėra suvokiamas kaip mažiau autentiškas nei žmogaus sukurtas turinys. Požiūris į prekių ženklą taip pat reikšmingai nesikeičia net jei turinio autorystė yra atskleidžiama. Teksto emocingumo lygis (žemas, vidutinis, aukštas) neturėjo reikšmingos įtakos vartotojų požiūriui.
Zhang & Gosline, 2023	Žmogaus, DI ar žmogaus ir DI sukurtas turinys reklamoje	Tyrimo tikslai: 1. Iširti, kaip žmonės vertina turinį, sukurtą žmogaus, DI, ar žmogaus ir DI bendradarbiaujant. 2. Nustatyti, ar žmonės rodo favoritizmą žmogui ar vengia DI. Hipotezės: (H ₁) Žmonės geriau vertins turinį, pažymėtą kaip sukurtą žmogaus; (H ₂) DI sukurtas turinys, jei nepažymėtas, bus įvertintas teigiamai; (H ₃) Žmogaus favoritizmas lemia šį šališkumą, o ne neigiamas požiūris į DI.	DI kurtas turinys buvo vertinamas teigiamai, kai autorystė nebuvo nurodyta. Dalyviai teikė pirmenybę turiniui, pažymėtam kaip žmogaus kurtam, net jei DI turinys buvo objektyviai geresnis. Nepažymėtas DI turinys gavo aukštesnius įvertinimus nei pažymėtas. Favoritizmas žmogaus naudai buvo labiau susijęs su socialinėmis normomis nei su DI vengimu.
Shantatula et al., 2024	DI sugeneruotas turinys (AIGC) ir vartotojų sugeneruotas turinys (UGC) socialinės medijos rinkodaroje	Iširti, kaip DI sugeneruotas turinys ir vartotojų sugeneruotas turinys veikia vartotojus, identifikuoti jų privalumus, trūkumus. Hipotezės: (H ₁) DI generuotas turinys turi reikšmingą poveikį socialinės medijos rinkodaros strategijai. (H ₂)	Dirbtinio intelekto naudojimas leidžia personalizuoti ir greitai sukurti didelį kiekį turinio, tačiau, vartotojų požiūriu, toks turinys stokoja autentiškumo ir kūrybiškumo. Vartotojų sugeneruotas turinys yra suprantamas

		Vartotojų sugeneruotas turi reikšmingą poveikį socialinės medijos rinkodaros strategijai.	kaip patikimesnis ir įtikinamesnis. Autoriai rekomenduoja įsivertinti abiejų strategijų naudojimo privalumus ir trūkumus, mat tai gali stiprinti vartotojų įsitraukimą ir lojalumą.
Lim & Schmalzle, 2024	Autorystės atskleidimo poveikis DI sukurtų ir žmogaus sukurtų sveikatos komunikacijos pranešimų vertinimui	Tyrimo tikslai: 1. Nustatyti, kaip autorystės atskleidimas veikia pranešimų vertinimą. 2. Ištirti, kaip žmonių nuostatos dėl DI moderoja šį efektą. Hipotezės: (H ₁) Dalyviai žemesniais balais vertins DI kurtus pranešimus, kai bus nurodyta, kad juos kūrė DI; (H ₂) Dalyviai pirmenybę teiks žmogaus kurtam turiniui, jei žinos autorių; (H ₃) Neigiama nuostata dėl DI moderos autorystės atskleidimo įtaką pranešimų vertinimui; (H ₄) Neigiama nuostata dėl DI moderos autorystės atskleidimo poveikį pasirinkimui.	Pirmame tyrime nustatyta, kad dalyviai vertino DI sukurtus pranešimus geriau, kai autorystė nebuvo atskleista. Atskleidus autorystę, reikšmingo skirtumo tarp žmogaus ir DI kurtų pranešimų nebuvo. Antrojo tyrimo metu neigiama nuostata dėl DI modero autorystės poveikį: neigiamas požiūris į DI buvo susijęs su didesniu DI kurtų pranešimų efektyvumo vertinimu.
Gu et al., 2024	DI įrankių sukurtų reklamų savybių (tikroviškumo, gyvybingumo, kūrybiškumo ir sintetiškumo) poveikis vartotojų suvokimui ir priimtinumui.	Nustatyti, kaip DI sukurtų reklamų savybės (tikroviškumas, gyvybingumas, kūrybiškumas ir sintetiškumas) veikia vartotojų suvokimą apie reklamos keistumą, intelektą ir priimtinumą. Hipotezės: (H ₁) Tikroviškumas mažina suvokiamą keistumą ir didina intelektą; (H ₂) Gyvybingumas mažina suvokiamą keistumą ir didina intelektą; (H ₃) Kūrybiškumas mažina suvokiamą keistumą ir didina intelektą; (H ₄) Sintetiškumas didina suvokiamą keistumą ir mažina intelektą; (H ₅) Suvokiamas keistumas mažina norą priimti reklamas; (H ₆) Suvokiamas intelektas didina reklamos priimtinumą.	Tikroviška, suvokiama kaip gyvybinga / energinga ir kūrybinga reklama didina vartotojų pasitikėjimą DI ir skatina palankumą matomai reklamai. Sintetiškumo jausmas sukelia nejaukumą, kuris mažina vartotojų norą priimti reklamą.
Park et al., 2024	„Instagram“ naudotojų požiūris į DI, nuomonės formuotojų ir paprastų vartotojų paskyras bei jose skelbiamą turinį	Nustatyti, kaip vartotojai suvokia DI sukurtą turinį, ar atpažįsta jį. Tyrimo klausimai: 1. Ar naudotojų prioritetai skiriasi pagal paskyros tipą? 2. Kaip skiriasi turinio kokybės vertinimai? 3. Ar naudotojai gali atpažinti DI turinį?	Respondentai dažnai negalėjo atskirti DI sukurtą turinį nuo žmogaus kuriamo, tačiau nuostatos apie turinio kokybę ir patikimumą skyrėsi priklausomai nuo paskyros tipo. Nuomonės formuotojų paskyros buvo vertinamos kaip patikimiausios ir profesionaliausios, tačiau mažiau autentiškos nei paprastų vartotojų. DI paskyros buvo vertinamos kaip inovatyvios, bet jomis stoka pasitikėjimo.
Lee et al., 2025	DI sukurtų sporto reklamų vertinimas: kaip reklamos modelio tipas, DI autorystės atskleidimo momentas ir šaltinio	Tyrimas siekė išsiaiškinti, kaip skirtingas DI autorystės atskleidimo momentas (prieš / po / nežinoma), reklamos modelio tipas (žmogus, virtualus modelis, skaitmeninis dvynys) ir šaltinio – žinutės atitiktis	Tyrimas parodė, kad reklamos, kuriose DI autorystė buvo atskleidžiama po reklamos peržiūros, buvo vertinamos palankiausiai, o vertinimai buvo ženkliai žemesni, kai

	bei žinutės (ne)atitiktis veikia vartotojų požiūrį į reklamą	paveikia vartotojų reklamos vertinimą. Hipotezės: (H ₁) Reklamos vertinimas priklauso nuo DI autorystės atskleidimo laiko; (H ₂) Vertinimas priklauso nuo modelio tipo (žmogus, virtualus, skaitmeninis dvynys); (H ₃) Šaltinio ir žinutės neatitiktis neigiamai veikia reklamos vertinimą; (H ₄) DI autorystės atskleidimo momentas ir modelio tipas interaktyviai veikia vertinimą esant neatitiktčiai.	autorystė buvo atskleidžiama prieš peržiūrą.
--	--	---	--

Išsami mokslinių šaltinių, aprašančių vartotojų požiūrį į DI valdomų sprendimų sukurtą rinkodaros turinį, apžvalga, apimanti tyrimų autorius, tyrimo objektus, įvardijamas teorijas, iškeltus tikslus ar hipotezes, naudotus metodus, gautus rezultatus bei tolimesnių tyrimų rekomendacijas, yra pateikta 1 priede.

Remiantis įvairių mokslinių tyrimų, nagrinėjančių vartotojų požiūrį į DI valdomų sprendimų sukurtą rinkodaros turinį, apžvalga, pastebima, kad DI taikymas rinkodaroje sukelia skirtingus vartotojų atsakus. Vieni tyrimai rodo, kad vartotojai DI sukurtą rinkodaros turinį vertina taip pat, kaip ir žmogaus sukurtą turinį, nurodant, kad DI turinio autorystės žymėjimas neturi reikšmingo poveikio vartotojų atsakui (Kirkby et al. 2023). Tačiau kiti tyrimai atskleidžia, kad, nepaisant to, jog DI pasižymi turinio generavimo sparta bei personalizacijos galimybėmis, toks turinys vartotojų yra vertinamas prasčiau jei jo autorystė yra atskleidžiama (Zhang & Gosline, 2023; Lim & Schmalzle, 2024, Shantatula et al., 2024; Lee et al., 2025).

Visgi, vartotojų požiūrio į DI sugeneruotą rinkodaros turinį socialiniuose tinkluose tema atskleidžia reikšmingą tyrimų spragą, susijusią su moksliniu turiniu rinkodaroje. Apžvelgti tyrimai (žr. 1 lentelė) orientuojasi į skirtingus aspektus, susijusius su vartotojų požiūriu į DI naudojimą rinkodaros turinio kūrimui. Kai kurie iš jų siekė išsiaiškinti tiesioginį vartotojų požiūrį į DI sukurtą turinį, tuo tarpu kiti analizavo veiksnius, galinčius formuoti šį požiūrį, kaip turinio charakteristikos (Argan et al., 2022; Park et al., 2024; Gu et al., 2024) ar pačių respondentų egzistuojantys požiūriai į DI ir robotus (Wu et al., 2021; Wu & Wen, 2021). Vis dėlto, šie tyrimai apima rinkodaros turinį, susijusį su produktais ar paslaugomis, kurios neturi tiesioginio ryšio su DI. Tai lemia įžvalgų apie DI technologijų vaidmenį kuriant specializuotą, su mokslu susijusį turinį, stoką, mat nėra aišku, kaip vartotojų požiūris ir reakcijos keičiasi, kai DI sukurtas turinys yra susijęs su inovatyviais DI taikymo produktų ar paslaugų kūrimo aspektais. Pavyzdžiui, kai DI sukurtos ir pažymėtos rinkodaros žinutės aprašo DI technologijų pagrindu sukurtus produktus ar paslaugas. Nėra žinoma, ar toks turinys yra vertinamas palankiau dėl aiškos turinio šaltinio ir reklamuojamų sprendimų sąsajos nei DI sukurtos žinutės apie su DI nesusijusius produktus ar paslaugas. Tai yra aktualu siekiant suprasti, ar vartotojai labiau pasitiki DI kuriamu ir pažymėtu turiniu, kai jis tiesiogiai atspindi technologinį kontekstą, ar priešingai – tokio turinio pritaikymas platesniam produktų ir paslaugų spektrui gali turėti neigiamą poveikį vartotojų atsakui. Šis tyrimas galėtų papildyti teorines žinias apie vartotojų elgseną, susijusią su DI ir suteikti praktinių įžvalgų specialistams, siekiantiems kurti veiksmingas rinkodaros strategijas, grįstas DI sukurtu turiniu socialiniuose tinkluose.

Be to, daugelis analizuotų tyrimų rėmėsi apklausos metodais (Wu & Wen, 2021; Gu et al., 2024; Shantatula et al., 2024; Park et al., 2024), kurie apribojo galimybes tyrėjams detaliau nagrinėti vartotojų sąveikas su konkrečiais DI sugeneruoto turinio tipais ir jų tarpusavio sąveiką. Nors tokie metodai leido numatyti bendrąsias tendencijas, jie nesuteikė išsamių įžvalgų apie tai, kaip vartotojai reaguoja į tekstą ir vizualus, kai jie pateikiami kartu. Tuo tarpu apžvelgti eksperimentiniai tyrimai, nors turintys potencialo užpildyti šią spragą, paprastai apsiribojo tik vieno elemento analize – tekstiniu turiniu (Kirkby et al., 2023), neatsižvelgiant į socialinių tinklų kontekstą, kuris įprastai apima atitinkamą vizualinę struktūrą. Atsižvelgiant į tai, kad socialiniuose tinkluose DI sugeneruotas rinkodaros turinys gali būti pateikiamas kaip kompleksinė tekstų ir vizualų visuma, būtina analizuoti šių elementų sąveiką ir jų poveikį vartotojų atsakui. Kadangi, dėl DI autorystę pažyminčių etikečių naujumo, iki šiol nėra aišku, ar vartotojai DI autorystę priskiria tik vienam iš šių turinio tipų, ar abu elementus suvokia kaip vienodai susijusius su DI naudojimu, šios tyrimų spragos užpildymas galėtų reikšmingai prisidėti prie vartotojų požiūrio ir įsitraukimo supratimo ir efektyvesnių DI sugeneruoto turinio strategijų formavimo socialiniuose tinkluose.

2. Dirbtinio intelekto sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose teorinis pagrindimas

2.1. Dirbtinio intelekto sukurto rinkodaros turinio socialiniuose tinkluose ypatybės

Socialiniai tinklai (angl. *social network sites*) gali būti apibrėžiami kaip internetinės platformos, kurios leidžia vartotojams susikurti viešą ar pusiau viešą profilį, užmegzti ryšius su kitais vartotojais ir matyti savo bei kitų vartotojų profilių sąrašą vienoje sistemoje (Boyd & Ellison, 2007). Šių tinklų esmė – socialinių ryšių kūrimas ir atvaizdavimas, todėl jie pasižymi stipria bendruomenine struktūra ir vartotojų tarpusavio sąveika. Tuo tarpu socialinė medija (angl. *social media*) yra platesnė sąvoka, apimanti mobiliąsias ir internetines platformas, leidžiančias vartotojams kurti, dalytis turiniu ir žymėti jį geografinėmis žymėmis, pasiekiant itin didelę auditoriją (Ouiridi et al., 2014). Tai apima įvairius komunikacijos formatus, tokius kaip tekstas, vaizdai, vaizdo įrašai, garso įrašai ir net žaidimai. Anot mokslinių šaltinių, socialinė medija iš esmės skiriasi nuo tradicinių masinių komunikacijos kanalų, nes skatina interaktyvią ir abipusę komunikaciją tarp naudotojų (Ouiridi et al., 2014; Geissinger, & Laurell, 2016).

Mokslinėje literatūroje galima rasti įvairių šių sąvokų klasifikacijų – kai kurie autoriai jas vartoja apibrėžiant atskiras kategorijas, o kiti jas sieja hierarchiniu ryšiu. Pavyzdžiui, Kaplan'as ir Haenlein'as (2010) socialinę mediją apibrėžia kaip internetinių taikomųjų programų visumą, paremtą „Web 2.0“ principais, kuri leidžia vartotojams kurti ir dalytis naudotojų sukurtu turiniu, ir šioje klasifikacijoje socialinius tinklus priskiria kaip vieną iš socialinės medijos tipų. Dėl šių terminų panašumo bei dalinio reikšmės sutapimo – abu apibrėžia skaitmeninės komunikacijos, turinio kūrimo ir bendruomeniškumo aspektus – mokslinėje literatūroje jie taip pat neretai vartojami kaip lygiaverčiai (Ouiridi et al., 2014; Geissinger & Laurell, 2016; Trakimavičiūtė, 2017). Todėl, siekiant užtikrinti koncepcinį nuoseklumą ir aprėpti platesnį aktualių šaltinių spektrą, šiame projekte sąvokos socialiniai tinklai ir socialinė medija vartojamos kaip sinonimiškos.

DI technologijų taikymas socialinių tinklų platformose nėra naujas reiškinys – jos jau kurį laiką integruojamos į įvairius sistemos veikimo aspektus, nuo turinio filtravimo ir personalizavimo iki reklamos srautų valdymo. Socialinių tinklų platformos, tokios kaip „Facebook“ ir „TikTok“, plačiai naudoja DI sprendimus, siekdamos valdyti turinį, personalizuoti reklamą ar moderuoti vartotojų skelbiamą problematišką turinį (Grandinetti, 2023). „Facebook“ platformoje DI technologijos daugiausiai naudojamos problematiško turinio moderavimui, į kurio procesą įsitraukia tiek žmogus, tiek mašininio mokymosi algoritmai, leidžiantys automatizuotai atpažinti ir filtruoti įžeidžiamą ar netinkamą turinį. Taip pat „Facebook“ naudoja DI algoritmais, siekiant rekomenduoti atitinkamą reklaminį turinį galimai tuo susidomėsiąsiais vartotojais. „TikTok“ išsiskiria personalizavimo galimybėmis – platformoje įdiegtos DI technologijos analizuoja vartotojų elgseną (peržiūrų laiką, paspaudimus, „patinka“ paspaudimus ir komentarus), kategorizuoja skelbiamą turinį į temas ir teikia individualizuotas rekomendacijas vartotojams, siekiant maksimizuoti vartotojų įsitraukimą (Kang & Lou, 2022; Grandinetti, 2023). Be to, DI įrankiai padeda kuriant turinį, siūlydami automatinį vaizdo redagavimą, filtrus, muzikos parinkimą ir kitus kūrybinius sprendimus, kurie leidžia net ir mažai patirties turintiems vartotojams lengvai kurti patrauklius vaizdo įrašus (Kang & Lou, 2022).

Vartotojų požiūris į tokį DI taikymą yra skirtingas. Kang'o ir Lou'o (2022) tyrimas parodė, kad dauguma „TikTok“ naudotojų yra palankiai nusiteikę DI atžvilgiu. Jie teigiamai vertina personalizuotą turinį ir patogumą, kurį suteikia socialiniame tinkle realizuotos DI grįstos

rekomendacijos. Respondentai pripažino, kad algoritmai labai greitai prisitaiko prie jų interesų ir pateikia aktualų turinį, kuris skatina ilgesnį naudojimąsi platforma. Tačiau kartu buvo pastebėta tam tikra dvejopa reakcija – nors personalizavimas laikomas patraukliu, vartotojai taip pat išreiškė susirūpinimą dėl autonomijos ir savo duomenų privatumo. Kai kurie respondentai jautė, kad jų turinio vartojimo įpročiai yra pernelyg stipriai kontroliuojami DI, o tai gali kelti abejonių dėl jų galimybės savarankiškai atrasti naują turinį.

Tokios vartotojų reakcijos rodo, kad jų patirtį DI grįstose socialinių tinklų platformose lemia ne vien technologijos veikimo principai, bet ir tai, kokią vertę vartotojams suteikia šiose platformose prieinamas turinys. Kadangi socialiniai tinklai iš esmės yra grindžiami turinio kūrimu, dalijimusi ir sąveika tarp vartotojų (Boyd & Ellison, 2007; Kaplan & Haenlein, 2010), turinys tampa centriniu šių platformų elementu. Socialiniuose tinkluose šis turinys dažniausiai pasireiškia dviem pagrindinėmis formomis – tekstu ir vizualiniais elementais – kurie kartu formuoja bendrą komunikacijos visumą ir daro reikšmingą įtaką vartotojų įsitraukimui.

2.1.1. Dirbtinio intelekto sukurtas tekstinis rinkodaros turinys socialiniuose tinkluose

DI sugeneruotas turinys socialiniuose tinkluose ne visuomet sulaukia vienareikšmiškos vartotojų reakcijos. Tyrime, kurį atliko Yang'as (2024), buvo analizuojama, kaip DI sukurti naujienų antraščių pavadinimai veikia vartotojų įsitraukimą socialinių tinklų platformoje „RED“ – buvo siekiama nustatyti, ar DI sugeneruotos antraštės gali paskatinti vartotojų įsitraukimą taip pat veiksmingai, kaip žmonių parašytos, bei kokie veiksniai lemia šį skirtumą. Tyrimas rėmėsi penkių etapų antraščių optimizavimo metodu, kuris įtraukė populiarių raktinių žodžių naudojimą, auditorijos poreikių analizę ir sukurtą turinio peržiūrą, kurią atliko žmogus, siekiant užtikrinti, kad antraštės būtų aiškos, patrauklios ir atitinkančios turinio kontekstą. Nustatyta, kad didesnio vartotojų įsitraukimo sulaukė tos antraštės, kurios siekė pritraukti vartotojų dėmesį kontrastingais kalbos elementais bei siekė sukelti smalsumą. Vis dėlto DI sugeneruotos antraštės turėjo beveik keturis procentais mažesnį paspaudimų (angl. *click-through rate*) rodiklį nei žmonių sukurtos (Yang, 2024).

Tyrimas, atliktas Ananthakrishnan'o ir Arunachalam'o (2022), nagrinėjo socialinių tinklų vartotojų požiūrį į DI sugeneruotas ir žmogaus sukurtas prekių ženklo reklamas, neatskleidžiant, kuris turinys sukurtas DI, siekiant išvengti šališko vertinimo. Respondentai vertino reklaminius pranešimus apie prabangų automobilį pagal turinio kūrybiškumą, tipą ir atitikimą prekių ženklui. Rezultatai atskleidė, kad net 69 procentai respondentų teikė pirmenybę DI sugeneruotam turiniui, o tai, anot autorių, rodo, kad vartotojai ne visada atskiria turinio šaltinį vien iš sukurtą turinio ir kad DI sukurti reklaminiai pranešimai gali būti veiksmingi. Kūrybiškumo vertinimo atžvilgiu DI sukurtas turinys buvo laikomas modernesniu ir įsimintinesniu, tačiau žmogaus sukurtas turinys labiau išsiskyrė savo emociniu poveikiu ir atitikimu prekių ženklui. Vertinant turinio tipą, žmogaus kurti pranešimai buvo laikomi labiau pramoginiais, o DI sugeneruoti – aiškesniais ir informatyvesniais (Ananthakrishnan & Arunachalam, 2022). Remiantis šiais tyrimų duomenimis, galima daryti prielaidą, jog DI pasižymi gebėjimu veiksmingai generuoti informacinį turinį socialiniams tinklams.

Būtent informatyvumas, kaip viena iš esminių turinio charakteristikų, glaudžiai siejasi su socialinių tinklų turinio tipologija, kuri dažniausiai išskiriama į informacinį, emocinį ir transakcinį turinį (Dolan et al., 2019; Shahbaznezhad et al., 2021; Wahid et al., 2023). Šis skirstymas remiasi unikaliomis turinio ypatybėmis, kurios daro skirtingą poveikį vartotojų elgsenai bei įsitraukimui:

- Informacinis turinys, dar vadinamas racionaliū arba edukaciniu, yra orientuotas į vartotojams naudą teikiančią ir objektyvią informaciją (Laskey et al., 1989). Toks turinys socialiniuose tinkluose gali apimti prekės ar paslaugos aprašymus, nuotraukas, įvykių informaciją ar įvairias edukacines žinutes (Dolan et al., 2019). Toks turinys dažniausiai susijęs su loginiu informacijos apdorojimu ir vartotojų sprendimais, kurie grindžiami faktais. Moksliniai tyrimai rodo, jog šio tipo turinys dažnai generuoja pasyvų vartotojų įsitraukimą, pavyzdžiui, „patinka“ reakcijas ir retai paskatina aktyvesnius veiksmus, tokius kaip komentarai ar dalijimasis (Cvijikj & Michahelles, 2013; Lee et al., 2018; Dolan et al., 2019). Nepaisant to, informacinis turinys yra itin svarbus, nes stiprina prekių ženklo autoritetą ir užtikrina jo patikimumą vartotojų akyse (Shahbaznezhad et al., 2021).
- Emocinis turinys išsiskiria tuo, kad jo tikslas – sužadinti vartotojų jausmus ir emocijas. Šis turinys gali apimti humoristinius tekstus, paveikslėlius ar vaizdo įrašus, įkvepiančias istorijas ar pramogines žinutes, kurios ne tik padeda sukurti stipresnį emocinį ryšį su auditorija, bet ir skatina bendravimą tarp pačių vartotojų (Dolan et al., 2019). Tyrimai atskleidžia, kad emocinis turinys dažnai sulaukia didesnio įsitraukimo, tačiau pasyvių reakcijų, tokių kaip „patinka“ paspaudimai (Lee et al., 2018; Dolan et al., 2019; Shahbaznezhad et al., 2021).
- Transakcinis turinys orientuojasi į tiesioginius veiksmus, skatinančius vartotojų įsitraukimą per nuolaidas, konkursus ar akcijas (Dolan et al., 2019). Šio turinio pagrindinis tikslas yra paskatinti vartotojus atlikti konkrečius veiksmus, tokius kaip pirkti prekes, registruotis ar dalyvauti lojalumo programose. Nors transakcinis turinys gali būti efektyvus skatinant pasyvų įsitraukimą, jo poveikis aktyviam vartotojų elgesiui, pavyzdžiui, komentarams, priklauso nuo platformos bei informacijos pateikimo būdo (Shahbaznezhad et al., 2021; Lee et al., 2018).

Naujausia apklausa, atlikta Jungtinėse Valstijose 2024 m. gegužę, parodė, kad emocinio turinio tipas yra labiausiai vertinamas vartotojų – net 50 proc. respondentų nurodė, kad jie teikia pirmenybę humoro elementus turinčioms reklamoms (Statista Research Department, 2024). Be humoro, taip pat populiarius yra ir informacinis turinys – 40 proc. vartotojų teigiamai vertina „kaip tai padaryti“ (angl. *how to*) pobūdžio reklamas, nes jos suteikia naudingos ir praktiškai pritaikomos informacijos (Statista Research Department, 2024). Tai patvirtina ir moksliniai tyrimai – Hassan'o (2024) tyrimo rezultatai atskleidė, kad informatyvumas, interaktyvumas, kūrybiškumas ir reklamos aktualumas yra pagrindiniai veiksniai, lemiantys teigiamą požiūrį į reklamą. Informatyvumas laikomas vienu svarbiausių aspektų, nes vartotojai socialiniuose tinkluose dažnai ieško reklamos, kuri teikia aktualią ir naudingą informaciją apie produktus ar paslaugas. Tiksliai, savalaikė ir išsami informacija padeda vartotojams priimti geresnius sprendimus dėl pirkimo, o tai skatina palankų požiūrį į reklamą (Hassan, 2024). Interaktyvumas yra dar vienas svarbus aspektas, nes vartotojai vertina galimybę tiesiogiai bendrauti su prekių ženklais, dalytis atsiliepimais ir gauti atsakymus greitai. Tyrimo duomenys parodė, kad didesnis reklamos interaktyvumas socialiniuose tinkluose skatina vartotojų įsitraukimą, o tai savo ruožtu teigiamai veikia jų požiūrį į prekių ženklą (Hassan, 2024). Kūrybiškumas taip pat buvo reikšmingas veiksnys – vartotojai labiau įsitraukia į išskirtinę, originalią ir netikėtą reklamą. Tyrimas parodė, kad kūrybiškas turinys labiausiai skatina teigiamą požiūrį į reklamą, o tai padidina prekių ženklo žinomumą ir potencialių klientų susidomėjimą (Hassan, 2024). Reklamos aktualumas – tai gebėjimas pritaikyti reklamą pagal vartotojų poreikius, interesus ir gyvenimo būdą. Tyrimo rezultatai patvirtino, kad reklama, kuri tiesiogiai atitinka vartotojų asmeninius poreikius, sukelia palankesnę požiūrį ir didina tikimybę, kad vartotojai įsitrauks į prekių ženklo komunikaciją (Hassan, 2024).

Atsižvelgiant į mokslinių tyrimų bei vartotojų apklausų rezultatus, galima teigti, kad DI geba efektyviai generuoti informacinio pobūdžio turinį (Ananthakrishnan & Arunachalam, 2022; Yang, 2024), kuris yra viena svarbiausių vartotojų vertinamų socialinių tinklų reklamos savybių (Statista Research Department, 2024; Hassan, 2024). Kadangi informatyvumas yra vienas pagrindinių veiksnių, lemiančių teigiamą vartotojų požiūrį į reklamą, DI generuojamos tekstinės žinutės gali sėkmingai prisidėti prie vartotojų įsitraukimo didinimo socialinių tinklų platformose (Yang, 2024). Informacinio tipas savo funkcija iš esmės atliepia moksliniam turiniui būdingą tikslą – perteikti aiškią, aktualią ir objektyvią informaciją, todėl yra itin tinkamas komunikacijai apie inovatyvius, su pačiu DI taikymu susijusius produktus ar paslaugas.

2.1.2. Dirbtinio intelekto sukurtas vizualinis rinkodaros turinys socialiniuose tinkluose

Socialinių tinklų aplinkoje informacija retai pateikiama vien tekstiniu pavidalu – dažniausiai ji formuojama kaip tekstinių ir vizualinių elementų visuma. Todėl, nepaisant DI gebėjimų generuoti tekstą, svarbu atsižvelgti ir į vizualinę raišką, kuri, veikdama kartu su tekstu, gali reikšmingai paveikti vartotojų suvokimą, įsitraukimą bei vertinimą.

Nors rinkodaros praktikoje dažnai pasitelkiami įvairūs reklamos formatai, apimantys tiek tekstinius pranešimus, tiek vizualinius elementus, moksliniai tyrimai pabrėžia, kad vartotojai itin palankiai vertina abiejų formatų derinį, kuris sustiprina reklamos efektyvumą. Kim'o et al. (2016) tyrime buvo analizuojama, kaip tekstiniai ir vizualiniai elementai veikia reklamos efektyvumą socialiai atsakingos komunikacijos kampanijose restoranų sektoriuje. Pagrindinis dėmesys buvo skiriamas vartotojų reakcijoms į skirtingus reklamos formatus – vien tik tekstinius pranešimus ir tekstinius pranešimus, papildytus vizualiniais elementais, bei kaip šie formatai veikia požiūrį į prekių ženklą, reklamą ir elgsenos ketinimus. Rezultatai parodė, kad reklamos, kurios derina tekstą su vizualiniais elementais, yra kur kas veiksmingesnės nei vien tekstinės žinutės. Dalyviai, peržiūrėję reklamas su paveikslėliais, pateikė pozityvesnius vertinimus apie reklamos patrauklumą, jos ryšį su socialinėmis iniciatyvomis bei prekių ženklo įvaizdį. Vizualų naudojimas reklamose ne tik suteikė daugiau informacijos vartotojams, bet ir palengvino jų kognityvinius procesus, todėl reklamos su vizualiniais elementais sukėlė didesnę pasitenkinimą ir stipriau motyvavo veikti (Kim et al., 2016).

Nors vartotojai reklamos turinį dažniausiai interpretuoja kaip vieningą visumą, vizualiniai elementai turi didesnę įtaką formuojant nuostatas nei tekstiniai elementai. Tai patvirtina vadinamasis „vizualinio pranašumo efektas“ (angl. *visual superiority effect*), kuris nurodo, jog vizualiniai reklamos elementai dažnai nustelbia tekstinius. Pavyzdžiui, Kim'o ir Kim'o (2022) tyrime, atvejais, kai reklamos turinys buvo ilgesnis ir sudėtingesnis, dalyviai labiau rėmėsi vaizdiniu turiniu, formuodami savo nuostatas. Tai reiškia, kad vizualiniai elementai, kurie yra lengviau ir greičiau apdorojami, turi didesnę poveikį vartotojų suvokimui nei tekstai, kurie reikalauja didesnių kognityvinių pastangų. Tyrimo metu buvo nustatyta, kad vartotojai, peržiūrėję reklamas su vizualiais elementais, labiau kreipė dėmesį į juos nei į tekstą. Remiantis Zihagh'o et al. (2024) analize, vartotojai dažniau interpretuoja tekstinius ir vizualinius elementus kaip vieningą visumą tada, kai jie atitinka vienas kitą pagal komunikacijos tikslą. Pavyzdžiui, jei tekstas ir vizualinis turinys perteikia vieningą žinutę, vartotojai greičiau įsitraukia į reklamą ir lengviau ją prisimena. Tačiau, jei vizualiniai ir tekstiniai elementai yra nesuderinti arba perteikia skirtingas žinutes, vartotojai dažniau juos suvokia kaip atskirus komponentus. Toks nesuderinamumas gali kelti kognityvinį disonansą, kai vartotojai stengiasi susieti skirtingas reklamos dalis. Pavyzdžiui, kai reklamoje naudojama neįprasta vizualika, bet pateikiama įprasta tekstinė informacija, vartotojai dažnai daugiau dėmesio skiria

neįprastam vizualiniam elementui, o tekstinis turinys lieka antraeilis. Tai rodo, jog vizualų naudojimas reklamos, kurioje egzistuoja nežymus neatitikimas tarp reklamos elementų, nukreipia vartotojų dėmesį nuo tekstinių žinučių ir tampa pagrindiniu vertinamu informacijos šaltiniu.

Nors ankstesni tyrimai pabrėžė, kad vizualiniai elementai daro didelę įtaką reklamos efektyvumui, ypač kai jie derinami su tekstu ir sudaro vieningą komunikacijos visumą (Kim et al., 2016; Kim & Kim, 2022), DI sugeneruotas turinys keičia tradicinius vizualinės reklamos suvokimo ir vertinimo mechanizmus. Exner'is et al. (2025) tyrime siekė išsiaiškinti, ar DI generuotos reklamos gali pasiekti arba pranokti žmogaus kurtų reklamų efektyvumą, bei identifikuoti pagrindinius vizualinius bruožus, kurie daro įtaką vartotojų suvokimui apie reklamos kilmę. Tyrimo rezultatai atskleidė, kad DI sukurti reklaminiai vaizdai gali būti tokie pat veiksmingi kaip žmogaus sukurti vaizdai, tačiau tik tais atvejais, kai vartotojai nesupranta, kad vaizdas yra sukurtas DI. Tyrimas parodė, kad DI generuotos reklamos pasiekia geresnį paspaudimų rodiklį (angl. *click-through rate*, *CTR*), kai jos „neatrodė kaip DI sukurtos“ – jei reklamos vizualai atrodė netikri, jie sulaukia mažesnio vartotojų įsitraukimo. Tyrėjai taip pat nustatė, kad tam tikri vizualiniai požymiai padeda vartotojams atskirti DI sukurtas reklamas nuo žmogaus kurtų. Pavyzdžiui, DI sukurti vaizdai dažnai yra labiau estetiški, mažesnės raiškos, pasižymi intensyvesnėmis spalvomis ir dideliu kontrastu. Taip pat, DI sukurtuose reklaminiuose vaizdiniuose rečiau matomas tekstas. Visgi, šie išskirti reklamų vaizdinių elementai taip pat gali lemti klaidinantį įspūdį apie matomas reklamas – nors DI įprastai susiduria su iššūkiais generuojant tekstą vaizduose, tais atvejais, kai tekstas yra tinkamai atkuriamas, vartotojai tokias reklamas neretai priskiria žmogaus sukurtam turiniui (Exner et al., 2025). Šie rezultatai parodo, kaip DI gali manipuluoti vartotojų suvokimu, naudodamas vizualinius elementus, kurie labiau primena žmogaus kurtą turinį, ir taip sumažinti savo „dirbtinumo“ įspūdį.

Grigsby'o et al. (2025) tyrime nagrinėjama, kaip DI sugeneruotų reklamų autorystės atskleidimas veikia vartotojų pasitikėjimą paslaugų teikėjais (angl. *source credibility*) ir jų požiūrį į reklamą. Tyrime buvo atlikti trys eksperimentai, kurie atskleidė, kad DI sugeneruotų reklamų autorystės atskleidimas dažniausiai lemia mažesnę vartotojų pasitikėjimą ir mažiau palankias reklamos vertinimo nuostatas. Tačiau tyrėjai taip pat nustatė, kad šį neigiamą efektą galima sušvelninti strategiškai derinant DI ir realaus pasaulio elementus. Pavyzdžiui, reklamos, kuriose DI įrankiai buvo naudojami tik reklamos fono elementams kurti, tačiau vizuale matomas pagrindinis subjektas, pavyzdžiui, paslaugų teikėjas, buvo realus, buvo vertinamos palankiau nei visiškai DI sugeneruotos reklamos. Be to, DI sugeneruotų reklamų su paslaugų teikėjų atvaizdais autorystės atskleidimas ypač sumažino pasitikėjimą paslaugų teikėju, nes paslaugų srityse, kur svarbus žmogiškasis ryšys, vartotojai yra jautresni informacijos šaltiniui. Šie rezultatai rodo, kad DI naudojimas reklamoje gali būti efektyvus, tačiau jo poveikis priklauso nuo to, kaip vizualiniai elementai yra pateikiami ir ar skaidriai komunikuojama apie DI naudojimą reklamos kūrime (Grigsby et al., 2025).

Socialinių tinklų rinkodaros praktikoje vizualinis turinys atlieka esminį vaidmenį siekiant padidinti vartotojų įsitraukimą ir reklamos efektyvumą. Tyrimai rodo, kad socialinių tinklų turinys, kuriame derinami tiek tekstiniai, tiek vizualiniai elementai, dažniausiai pasiekia aukštesnius įsitraukimo rodiklius nei vien tekstinis turinys (Fox et al., 2019; Wani & Ahmad, 2024). Socialiniai tinklai, tokie kaip „Instagram“ ar „TikTok“, yra sukurti vizualiniam turiniui publikuoti, todėl šiose platformose tekstiniai pranešimai negali būti skelbiami atskirai – jie privalo būti integruoti su vaizdiniais elementais, tokiais kaip nuotraukos ar vaizdo įrašai. Ši platformų architektūra lemia, kad vizualinis turinys tampa neatsiejama komunikacijos dalimi, formuojančia vartotojų patirtį ir įsitraukimą.

Wortel'o et al. (2024) tyrimas nagrinėjo DI sukurtų reklamų atskleidimo poveikį socialiniuose tinkluose, atsižvelgiant į tai, kad tokiose platformose kaip „Instagram“ vizualinis turinys yra būtinas komunikacijos elementas. Nors ankstesni tyrimai parodė, kad vartotojai paprastai labiau įsitraukia į reklamą, kai ji apjungia tekstinius ir vizualinius elementus (Fox et al., 2019; Wani & Ahmad, 2024), iki šiol nebuvo aišku, ar vartotojų reakcijos į DI sugeneruotą vizualinį turinį skiriasi nuo reakcijų į DI sukurtą tekstą. Atsižvelgdami į tai, kad vizualiniai elementai dažnai yra laikomi svarbesniais reklamos suvokimui nei tekstas (Kim & Kim, 2022), tyrėjai kėlė hipotezę, jog DI sugeneruoto vaizdo autorystės atskleidimas gali turėti stipresnį neigiamą poveikį reklamos vertinimui nei DI sukurto teksto autorystės atskleidimas. Eksperimente buvo tiriamos trys skirtingos sąlygos: reklamos be jokios informacijos apie DI naudojimą, reklamos su informacine žinute, kad reklamos tekstas buvo sukurtas DI ir reklamos su paaiškinimu, kad reklamos vizualas buvo sugeneruotas DI. Svarbu pabrėžti, jog tyrime informaciniai paaiškinimai apie DI naudojimą buvo pateikiami kaip dalis reklamos teksto, tam skirtame teksto laukelyje „Instagram“ platformoje – ne naudojantis naujausiomis DI autorystę atskleidžiančiomis etiketėmis. Rezultatai parodė, kad DI sugeneruotų reklamų autorystės atskleidimas reikšmingai sumažino vartotojų teigiamą požiūrį į reklamą, tačiau neturėjo tiesioginio poveikio nei prekių ženklo vertinimui, nei reklamos šaltinio patikimumui. Be to, nustatyta, kad DI atskleidimas mažino manipuliacinio ketinimo suvokimą – vartotojai, matydami aiškų DI atskleidimą, jautėsi mažiau manipuliuojami. Nepaisant to, bendras reklamos vertinimas vis tiek buvo žemesnis, patvirtinant hipotezę, jog DI sugeneruotas turinys gali sukelti neigiamas emocijas vien dėl paties fakto, kad reklama sukurta naudojant DI. Tyrime taip pat nebuvo nustatyta reikšmingų skirtumų tarp DI sukurto teksto ir vaizdo atskleidimų poveikio – tiek tekstinis, tiek vizualinis DI atskleidimas turėjo panašų poveikį vartotojų nuostatoms.

Tyrime taip pat analizuotas vartotojų algoritmų vengimo (angl. *algorithm aversion*) poveikis reklamos suvokimui. Rezultatai parodė, kad vartotojai, kurie skeptiškai vertina DI, prasčiau vertino prekių ženklą, kai sužinojo, jog reklama buvo sukurta naudojant DI. Tačiau šis neigiamas požiūris nepaveikė nei pačios reklamos vertinimo, nei jos šaltinio patikimumo. Tai reiškia, kad nors DI autorystės atskleidimas gali paskatinti neigiamą požiūrį į prekių ženklą, vertinant jį kaip mažiau „patrauklų“ ar „gerą“, jis neturi tiesioginio poveikio vartotojų požiūriui į pačią reklamą ar pasitikėjimui prekių ženklu. Šis rezultatas rodo, kad vartotojų skepticizmas DI atžvilgiu nėra vienodai nukreiptas į visus reklamos aspektus, o labiausiai susijęs su požiūriu į prekių ženklą. Šie rezultatai rodo, kad reklamos skaidrumas apie DI įrankių naudojimą kuriant tiek tekstą, tiek vizualus gali turėti dvigubą poveikį: nors jis mažina manipuliatyvaus ketinimo suvokimą, kartu jis gali pabloginti bendrą vartotojų požiūrį į prekių ženklą (Wortel et al., 2024).

Apibendrinant, vizualiniai elementai socialinių tinklų turinyje atlieka itin svarbų vaidmenį formuojant vartotojų atsaką. Tyrimai rodo, kad vizualinė raiška dažnai dominuoja prieš tekstinę informaciją, nes yra lengviau ir greičiau apdorojama, ypač dinamiškoje socialinių tinklų aplinkoje (Kim & Kim, 2022). Be to, vizualinių ir tekstinių komponentų integruotas formatas sustiprina reklamos įtaigumą, padidina vartotojų įsitraukimą bei palankiai veikia požiūrį į prekių ženklą (Kim, Kim, & Kim, 2016; Fox et al., 2019; Wani & Ahmad, 2024). Atsižvelgiant į šiuos aspektus, galima teigti, jog DI turinio efektyvumas socialiniuose tinkluose tiesiogiai priklauso nuo to, kaip sėkmingai turinio tipai, platformų specifika ir pateikimo formatai yra suderinti su vartotojų lūkesčiais bei elgsenos ypatybėmis. Nors DI sukurtas turinys gali būti toks pat efektyvus kaip žmogaus kurtas, jei vartotojai nežino jo kilmės (Exner et al., 2025), DI autorystės atskleidimas dažnai mažina reklamos vertinimą ir prekių ženklo patrauklumą (Grigsby et al., 2025; Wortel et al., 2024). Svarbu atkreipti

dėmesį į tai, kad ne visi vartotojai vienodai reaguoja į DI autorystės atskleidimą. Tyrimai rodo, kad vartotojai, turintys stipresnę algoritmų vengimo tendenciją, yra linkę skeptiškiau vertinti DI sugeneruotą turinį, nepaisant to ar tai tekstas, ar vizualai (Wortel et al., 2024). Tačiau ši skeptiška reakcija pirmiausia pasireiškia per neigiamą požiūrį į prekių ženklą, o ne į pačią reklamą ar jos šaltinio patikimumą.

2.2. Vartotojų atsaką į dirbtinio intelekto sukurto turinio autorystę paaiškinančios teorijos

Siekiant paaiškinti, kodėl vartotojai nevienodai reaguoja į DI sukurto turinio autorystės atskleidimą, būtina apžvelgti vartotojų elgseną paaiškinančius požiūrius. Vienas iš tokių požiūrių yra **algoritmų vengimo teorija** (angl. *algorithm aversion theory*), kuri paaiškina, kodėl vartotojai skeptiškai vertina DI priimamus sprendimus net tada, kai jie yra objektyviai tiksūs ir efektyvūs. Vartotojų atsakas į DI turinį taip pat gali būti paaiškintas remiantis **įtaigos atpažinimo modeliu** (angl. *persuasion knowledge model, PKM*), apibrėžiančius vartotojų skepticizmą atvejais, kai vartotojai atpažįsta reklamos šaltinio manipuliatyvumą.

Algoritmų vengimo teorija, pirmą kartą taip įvardinta Dietvorst'o et al. (2015) tyrime, remiasi idėja, kad vartotojai dažnai nepasitiki DI sprendimais, net jei jie objektyviai yra efektyvesni ar tikslesni nei žmonių priimami sprendimai (Dietvorst et al., 2015; Dietvorst et al., 2018; Castelo et al., 2019; Wortel et al., 2024; Haupt et al., 2025). Algoritmų vengimas gali priklausyti nuo atliekamos užduoties pobūdžio – žmonės labiau pasitiki DI priimtais sprendimais, kai jiems suformuluota užduotis yra objektyvi ir paremta skaičiavimais, tačiau skeptiškai vertina DI sprendimus, jei užduotis atrodo subjektyvi ir reikalaujanti žmogiškų savybių, tokių kaip intuicija ar emocinis supratimas (Castelo et al., 2019). Pavyzdžiui, vartotojai lengviau priima DI rekomendacijas dėl finansinių prognozių ar maršrutų optimizavimo, nes šios užduotys atrodo grįstos aiškiais skaičiavimais. Tyrimai rodo, kad vartotojai dažnai suvokia algoritmus kaip neturinčius pakankamų gebėjimų atlikti subjektyvias užduotis, pavyzdžiui, meno kūrimą ar emocijų atpažinimą, nepaisant to, kad DI gali generuoti tekstą (Pavlik, 2023), kurti vaizdus (OpenAI, n.d.), muziką (Civit et al., 2022). Tai rodo, kad algoritmų vengimas kyla ne tik iš objektyvių DI veikimo trūkumų, bet ir iš žmonių išankstinių nuostatų bei psichologinių mechanizmų.

Be to, Dietvorst'o et al. (2018) tyrimas atskleidė dar vieną svarbų algoritmų vengimo aspektą – netobulumo netoleranciją. Jų eksperimentas parodė, kad žmonės yra linkę atsisakyti net ir objektyviai tikslesnių algoritmų, jei pastebi, jog šie daro klaidų. Tai reiškia, kad nors DI prognozės gali būti geresnės nei žmonių, vienas klaidingas algoritmo sprendimas gali stipriai sumažinti pasitikėjimą juo. Tyrimo rezultatai rodo, kad žmonės yra atlaidesni žmonių klaidoms nei algoritmų klaidoms, net jei algoritmai klysta rečiau.

Algoritmų vengimas nėra pastovus – tam tikros intervencijos gali sumažinti šį efektą. Pavyzdžiui, vartotojai gali tapti atviresni DI sprendimams, jei jam formuojamos užduotys yra pateikiamos kaip objektyvios arba jei DI gebėjimai yra pristatomi kaip „žmogiškesni“ (Castelo et al., 2019). Taip pat, algoritmų vengimo efektą galima sumažinti arba neutralizuoti priskiriant kontrolę žmogui – jei žmonės gali keisti algoritmo sprendimus (Dietvorst et al., 2018, Haupt et al., 2025). Pavyzdžiui, kai DI sugeneruotas turinys pateikiamas kaip peržiūrėtas ir patvirtintas žmogaus, vartotojai nevertina jo prasčiau nei žmogaus sukurto turinio. Tai rodo, kad vartotojams svarbu žinoti, jog žmogus išlieka atsakingas už galutinį sprendimą, net jei DI atlieka didžiąją dalį techninio darbo.

Haupt'as et al. (2025) savo tyrime pabrėžia, kad algoritmų vengimo teorija yra ypač aktuali vartotojų reakcijoms į DI sugeneruotą turinį paaiškinti. Vartotojai skeptiškai žiūri į DI sukurtą turinį, kai yra aiškiai deklaruojama, jog turinys sukurtas be žmogaus įsikišimo. Tyrimai rodo, kad žmonės dažnai mano, jog DI trūksta žmogiškų savybių, tokių kaip empatija, moralinis supratimas ar kūrybiškumas, todėl jiems kyla abejonės dėl DI gebėjimo sukurti patikimą ir kokybišką turinį (Haupt et al., 2025).

Vartotojų požiūrį į DI turinį taip pat gali veikti jų įsitikinimai apie DI naudojimą. Minėtame tyrime (Haupt et al., 2025) nustatyta, jog tie asmenys, kurie laikė DI etišku įrankiu, buvo labiau linkę priimti DI sukurtą turinį, net jei jo kūrimo procesas nebuvo visiškai kontroliuojamas žmogaus. Tačiau tie asmenys, kurie jautė nepasitikėjimą ar laikė DI naudojimą problematišku, buvo kritiškesni ir skeptiškesni turinio patikimumo atžvilgiu, jei jį daugiausiai kūrė DI (Haupt et al., 2025). Šie rezultatai pagrindžia, kad vartotojų požiūris į DI sugeneruotą turinį priklauso ne tik nuo pačio turinio kokybės, bet ir nuo to, kiek jie jaučiasi užtikrinti, kad žmogus išlieka atsakingas už jo kokybę. Vartotojų neigiamas požiūris į DI dažnai kyla dėl baimės prarasti kontrolę arba dėl nepasitikėjimo DI gebėjimu priimti etiškai teisingus ir patikimus sprendimus. Tačiau jei DI naudojimas pristatomas kaip žmogaus prižiūrimas procesas, skeptiškumas sumažėja, o vartotojai yra linkę vertinti tokį turinį taip pat palankiai kaip ir žmogaus kurtą.

Remiantis aprašytais tyrimais, galima teigti, jog algoritmų vengimo teorija paaiškina, jog vartotojai gali skeptiškai vertinti DI sugeneruotą turinį dėl vartotojų polinkio mažiau pasitikėti technologiniais sprendimais, kai užduotys apima subjektyvius vertinimus, kūrybiškumą ar emocinį supratimą (Castelo et al., 2019; Dietvorst et al., 2015; Dietvorst et al., 2018). Be to, žmonės yra labiau linkę toleruoti žmogaus nei algoritmo klaidas, net jei algoritmas klysta rečiau (Dietvorst et al., 2018). Visgi vartotojų pasitikėjimas DI didėja, jei technologijos naudojimas pristatomas kaip žmogaus prižiūrimas procesas arba jei vartotojai turi galimybę išlaikyti kontrolę (Haupt et al., 2025). Tokiu atveju DI sugeneruotas turinys gali būti vertinamas panašiai palankiai kaip žmogaus kurtas. Taigi, vartotojų atsakas į DI sukurtą turinį socialiniuose tinkluose priklauso ne tik nuo objektyvių kokybės rodiklių, bet ir nuo jų nuostatų, pasitikėjimo lygio bei suvokiamo žmogaus vaidmens turinio kūrimo procese (Haupt et al., 2025; Wortel et al., 2024).

Įtaigos atpažinimo modelis, kurį sukūrė Friestad'as ir Wright'as (1994), aiškina, kaip vartotojai laikui bėgant išmoksta atpažinti įvairias įtikinimo taktikas ir kaip jie reaguoja į rinkodaros specialistų bandymus daryti jiems įtaką. Ankstesniuose šaltiniuose aprašytos tradicinės įtikinimo teorijos dažniausiai nagrinėjo, kaip žmonės priima pateikiamą informaciją ir kokį poveikį ji turi jų nuostatoms bei elgesiui. Tačiau šis Friestad'o ir Wright'o (1994) modelis išplečia šį požiūrį, pabrėždamas, kad vartotojai nėra vien pasyvūs reklamos ar pardavimo veiksmų stebėtojai bei aktyviai mokosi atpažinti įtikinimo strategijas, jas vertina ir pritaiko savo reakcijas priklausomai nuo asmeninių tikslų bei ankstesnės patirties. Modelis grindžiamas prielaida, kad įtaigos atpažinimas nėra įgimtas – jis yra formuojamas per asmenines ir socialines patirtis, stebint rinkodaros žinutes ir analizuojant kitų žmonių nuomonę apie reklamos bei pardavimo taktikas. Šis procesas nuolat kinta priklausomai nuo naujų technologijų, kultūrinių tendencijų bei asmeninių patirčių (Friestad & Wright, 1994).

Įtaigos atpažinimo modelis apibūdinamas per tris pagrindines žinių kategorijas:

- Įtaigos žinios (angl. *persuasion knowledge*) – vartotojų supratimas apie tai, kaip veikia įtikinimo procesai, kokios taktikos naudojamos reklamoje ar pardavimuose bei kaip šios taktikos gali paveikti jų sprendimus. Šios žinios leidžia vartotojams atpažinti situacijas,

kuriose naudojama įtaiga, ir atitinkamai reaguoti: nuo kritiško situacijos vertinimo iki visiško atmetimo ar priėmimo.

- Agentų žinios (angl. *agent knowledge*) – tai vartotojų turimos žinios apie įtikinėjimą vykdančius subjektus, tokius kaip reklamos agentūros, prekių ženklai ar organizacijos. Vartotojai vertina jų ketinimus, patikimumą ir kompetenciją, o šios išvados gali nulemti, kaip jie reaguoja į situacijas, kuriose pasitelkiama minėta įtaiga.
- Teminės žinios (angl. *topic knowledge*) – tai informacija apie konkretų produktą, paslaugą ar temą, kuri yra įtikinimo proceso centre. Kuo daugiau vartotojas žino apie atitinkamą temą, tuo kritiškiau jis gali vertinti reklamose ar kitose situacijose naudojamus teiginius ir mažiau priklausyti nuo įtaigos taktikų poveikio (Friestad & Wright, 1994).

Įtaigos atpažinimo modelis teigia, kad vartotojai įgyja įtikinimo žinias per gyvenimišką patirtį ir socialinę sąveiką. Pavyzdžiui, vaikai iš pradžių dažniausiai nesuvokia, kad reklamos tikslas yra paveikti jų pirkimo elgseną, tačiau su amžiumi jie išmoka atpažinti šį poveikį ir suformuoja strategijas, kaip su juo susidoroti (Friestad & Wright, 1994).

Kai vartotojai atpažįsta rinkodaroje naudojamas taktikas, jie gali reaguoti įvairiais būdais. Jei jie supranta, kad reklamos ar pardavimo strategijos siekia manipuliuoti jų elgseną, jie gali suformuoti gynybinius mechanizmus, kurie, siekiant išlaikyti kontrolę, padeda sąmoningai atsiriboti nuo reklamos poveikio. Vieni vartotojai gali tam pasitelkti savo fantaziją, kiti – ignoruoti atitinkamus „agento veiksmus“. Tiesa, autoriai šiam elgesiui priskiria ne tik veiksmus, kurių vartotojas imasi, bet ir jo kognityvinę elgesį, įskaitant galvojimą apie įtikinti bandančius subjektus, pačią situaciją (Friestad & Wright, 1994). Tačiau žinojimas apie galimą įtaigą ne visada sukelia neigiamą reakciją – kai kurios rinkodaros strategijos gali būti vertinamos kaip etiškos ir naudingos, jei vartotojai mano, kad jos pateikia svarbią ir aktualią informaciją. Pavyzdžiui, jei reklama yra aiškiai informatyvi ir atitinka vartotojo poreikius, jis gali ją suvokti kaip naudingą, o ne manipuliatyvią. Tokiais atvejais įtaigos žinios padeda vartotojams geriau įvertinti reklamos turinį ir priimti informuotus sprendimus.

Tyrimai rodo, kad įtaigos žinių modelis gali padėti paaiškinti, kaip vartotojai reaguoja į DI generuojamus vaizdus ir vaizdo įrašus reklamoje. Campbell'as et al. (2022) vartotojų atsaką aiškinamą per tris esminius veiksnius: reklamos įtikimumą (angl. *verisimilitude*), kūrybiškumą (angl. *creativity*) ir suvokiamą melagingumą (angl. *awareness of falsity*). Tyrimo rezultatai atskleidžia, kad didesnis vartotojų įsitikinimas reklamos autentiškumu koreliuoja su jos didesniu veiksmingumu. Tačiau, kai vartotojai suvokia reklamos sintetinę kilmę, jų skepticizmas ir nepasitikėjimas prekių ženklu gali sustiprėti, o tai gali turėti neigiamą poveikį reklamos efektyvumui.

Šiuos rezultatus patvirtina ir Powers'o et al. (2023) tyrimas, kuris analizavo, kaip vartotojai reaguoja į reklamas, kuriose naudojamos giliosios klastotės, ir kaip jų suvokimas apie reklamos patikimumą keičiasi, kai atskleidžiama, kad tai yra DI sukurtas turinys. Tyrimas atskleidė, kad vartotojai, kurie sužino, jog reklamoje naudojamas pseudoportretas yra sintetinė žmogaus kopija, dažnai tampa skeptiškesni, mažiau pasitiki prekių ženklu, o tai savo ruožtu mažina jų ketinimus pirkti. Įtaigos žinių modelis aiškina šią reakciją per įtaigos žinių aktyvumą: kai vartotojai supranta, kad įtikinimo bandymas yra grindžiamas manipuliatyvia technologija, jie gali pradėti taikyti gynybinius mechanizmus, tokius kaip skeptiškas reklamos vertinimas ar jos visiškas atmetimas. Tai rodo, kad DI sprendimais grįstos reklamos efektyvumas gali priklausyti nuo to, ar vartotojai suvokia jas kaip autentiškas ar manipuliatyvias (Campbell et al., 2022; Powers et al., 2023).

Apibendrinant galima teigti, jog vartotojų reakcijos į DI sugeneruotą turinį gali būti aiškinamos remiantis keliomis teorinėmis prieigomis, kurios pabrėžia tiek vartotojų motyvaciją naudotis technologijomis, tiek jų skepticizmą DI atžvilgiu. Algoritmų vengimo teorija rodo, kad žmonės yra linkę skeptiškai vertinti DI sprendimus, ypač jei jų atliekamos užduotys atrodo reikalaujančios žmogiškų savybių, tokių kaip emocinis supratimas ar kūrybiškumas (Dietvorst et al., 2018; Castelo et al., 2019). Tyrimai atskleidžia, kad net ir objektyviai tikslesni DI sprendimai gali būti atmetami vien dėl to, kad juos priėmė ne žmogus (Haupt et al., 2025). Be to, įtaigos atpažinimo modelis (Friestad & Wright, 1994) nurodo, kad vartotojai, supratę, jog DI kuriamas turinys yra reklamos ar manipuliavimo dalis, gali aktyvuoti gynybinius mechanizmus – didesnę skeptiškumą, nepasitikėjimą ar net tiesioginį tokio turinio atmetimą. Empiriniai tyrimai rodo, kad vartotojai, sužinoję, jog reklama ar kita informacija yra sugeneruota DI, dažnai tampa kritiškesni jos atžvilgiu, vertina ją kaip mažiau patikimą ir gali neigiamai vertinti patį prekių ženklą (Campbell et al., 2022; Powers et al., 2023).

Nors DI turinys gali tenkinti tam tikrus vartotojų poreikius, didžioji dalis teorinių modelių ir ankstesnių tyrimų rezultatų rodo, kad vartotojų reakcijos į DI generuojamą turinį dažnai yra išskirtinai neigiamos. Skepticizmas kyla tiek dėl išankstinių nuostatų prieš DI algoritmus, tiek dėl nepasitikėjimo jų gebėjimu perteikti autentišką, etišką ir emociškai reikšmingą informaciją. Taigi, vartotojų reakcijos į DI sukurtą turinį nėra vienareikšmiškos, tačiau dominuojančios tendencijos rodo didelį nepasitikėjimą ir kritiškumą šių technologijų naudojimo socialinių tinklų rinkodaroje atžvilgiu.

2.3. Vartotojų atsakas į dirbtinio intelekto sukurtą rinkodaros turinį ir jo autorystės žymėjimą

Nepaisant augančio DI sprendimų taikymo rinkodaroje, DI sukurtą turinį žyminčios etiketės dar tik imamos integruoti į rinkodaros praktikas. Šių etikečių naujumas lemia poreikį analizuoti vartotojų atsaką į DI sukurtą turinį bei jo autorystės žymėjimą, siekiant suprasti, kaip tai paveikia vartotojų pasitikėjimą, požiūrį į turinį bei išitraukimą. Visgi, siekiant plačiau pažvelgti į DI turinio žymėjimo socialiniuose tinkluose problematiką ir įgyti išsamesnių žinių apie vartotojų požiūrį į DI iš įvairių perspektyvų, tyrime pasitelkta platesnė DI taikymo analizė, neapsiribojant vien socialiniais tinklais.

Siekiant geriau suprasti vartotojų atsaką į DI sukurtą turinį bei jo autorystės atskleidimo poveikį vartotojų atsakui, žemiau (žr. 2 lentelė) pateikiami mokslinių tyrimų rezultatai, apžvelgiantys vartotojų atsakus skirtingose sąveikose su DI bei jų poveikį priimamiems sprendimams.

2 lentelė. Vartotojų atsakas į DI sukurtą turinį ir jo autorystės atskleidimo poveikis vartotojų sprendimams

Autorius (-iai), metai	Sprendimo priėmėjas	Sprendimo tipas	Vartotojų atsakas	Poveikis sprendimui
Luo et al., 2019	Vartotojas	Pirkti produktus po pokalbio su DI robotu	Neigiamas, kai tapatybė atskleista prieš pokalbį; neutralus ar teigiamas, kai atskleista po pokalbio	Mažiau pirkimų, kai tapatybė atskleista prieš pokalbį; daugiau pirkimų, kai po pokalbio
Powers et al., 2023	Vartotojas	Pirkti produktus ar pasitikėti informacijos šaltiniu	Sumažėjęs pasitikėjimas po DI naudojimo reklamoje atskleidimo	Sumažėję pirkimo ketinimai
Arango et al., 2023	Vartotojas	Aukoti labdarai po DI sukurto vaizdo peržiūros ir jo autorystės atskleidimo	Atskleidus DI naudojimą, sumažėjusi vartotojų jaučiama empatija	Mažėja aukojimo ketinimai

Baek et al., 2024	Vartotojas	Aukoti labdarai po DI sukurtos reklamos peržiūros ir jos autorystės atskleidimo	Reklamos vertinamos kaip mažiau patikimos, tačiau reakcija taip pat priklauso nuo DI „žmogiškumo“ suvokimo	Sumažėjusios aukojimo sumos, ypač jei DI laikomas „nežmogišku“
Bui et al., 2024	Vartotojas	Lankytis vietovėse po DI sukurtų vizualizacijų peržiūros ir jų autorystės atskleidimo	Mažesnis autentiškumo suvokimas DI sukurtoms vizualizacijoms	Mažesnis ketinimas lankytis, jei DI autorystė atskleista
Carney et al., 2024	Vartotojas	Įsitraukti į turinį, kurį sukūrė DI	Jaučiamas silpnas socialinis ryšys su turinio kūrėju	Mažesnis įsitraukimas
Chen et al., 2025	Vartotojas	Įsitraukti į trumpus vaizdo įrašus apie produktus, kurių sukūrimui naudotas DI	Turinio kokybė suvokiama kaip mažesnė, tačiau turinio autorystės žymėjimas sužadina smalsumą	Didesnis įsitraukimas

Detali mokslinių tyrimų analizė, kurioje apžvelgiamas vartotojų atsakas į DI sukurtą turinį ir jo žymėjimą, nurodant tyrimų autorius, objektą, taikytas teorines prieigas, suformuluotus tikslus ar hipotezes, naudotus metodus, pasiektus rezultatus bei rekomendacijos tolimesniems tyrimams, pateikiama 2 priede.

Luo‘o et al. (2019) tyrime buvo atliktas eksperimentas, skirtas ištirti, kaip pokalbių robotų tapatybės atskleidimas veikia klientų pirkimų rodiklius. Tyrime nagrinėtas tapatybės atskleidimo laiko poveikis: prieš pokalbį, po pokalbio arba po sprendimo pirkti priėmimo. Tyrimo rezultatai parodė, kad pokalbiai su tapatybės neatskleidžiančiais pokalbių robotais buvo tokie pat efektyvūs kaip ir su patyrusiais darbuotojais (žmonėmis), paskatinę pirkimą net 23,7 proc. atvejų. Visgi, pokalbių robotų tapatybės atskleidimas prieš pokalbį ženkliai sumažino pirkimų rodiklius – net 79,7 proc. Jei pokalbių robotų tapatybė buvo atskleista po pokalbio arba po sprendimo priėmimo, vartotojai pirkti buvo linkę labiau. Autoriai taip pat nustatė, kad klientai pokalbių robotus suvokė kaip mažiau empatiškus ir išmanančius, nors pokalbių įrašai parodė, kad jų kompetencija prilygo patyrusių darbuotojų gebėjimams. Todėl, remiantis šiuo tyrimu, galima teigti, jog subjektyvūs klientų įsitikinimai apie DI technologijas neigiamai veikia vartotojų sprendimą pirkti.

Powers‘o et al. (2023) tyrime buvo analizuojamas išmaniųjų vaizdo klastočių naudojimo atskleidimo poveikis vartotojų pirkimo ketinimams ir informacijos šaltinio patikimumo suvokimui. Anot tyrimo rezultatų, išmaniųjų vaizdo klastočių naudojimo reklamoje atskleidimas reikšmingai sumažino vartotojų, ypač moterų ir didesnes pajamas gaunančių asmenų, pirkimo ketinimus ir šaltinio patikimumo suvokimą. Vartotojams matomo pseudoportreto fizinio patrauklumo ir ekspertiškumo suvokimas reikšmingai nesikeitė.

Anot Arango‘o et al. (2023) DI sukurtų vaizdų naudojimas reklamoje turi neigiamą poveikį vartotojų jaučiamai empatijai. Autorių tyrime buvo nagrinėjamas vartotojų atsakas į DI sukurtus vaizdus, naudojamus labdaros reklamoje – kaip žinojimas apie vaizdo netikrumą veikia vartotojų emocijas, empatiją ir aukojimo ketinimus. Mokslininkai atliko tris eksperimentinius tyrimus: pirmame tyrime buvo nustatyta, kad kai žmonės žino, jog vaizdas sukurtas DI technologijų, jų empatija mažėja, o tai silpnina kaltės jausmą ir mažina aukojimo ketinimus; antrasis tyrimas nagrinėjo, kaip įvairūs motyvai, kuriais labdaros organizacijos gali pateisinti DI vaizdų naudojimą, keičia vartotojų reakcijas. Rezultatai parodė, kad etinių motyvų akcentavimas, pavyzdžiui, privatumo apsauga, gali sušvelninti neigiamą poveikį, tačiau instrumentinių motyvų, tokių kaip kaštų taupymas, akcentavimas

dažniausiai pablogina rezultatus. Trečiajame tyrime buvo analizuojama, ar ekstremalios aplinkybės, tokios kaip stichinės nelaimės, daro įtaką vartotojų požiūriui į DI vaizdus. Rezultatai parodė, kad nelaimių atvejais, kai realių vaizdų, planuojamų naudoti reklamose, gauti nėra įmanoma, DI sukurtų veidų naudojimas yra suvokiamas palankiau ir pasiekia panašius rezultatus, kaip ir realių vaizdų naudojimas. Apibendrinant, tyrimas atskleidė, kad DI sugeneruoti vaizdai nėra efektyvūs siekiant paskatinti vartotojų ketinimą aukoti bei neigiamai veikia jaučiamą empatiją, tačiau šis poveikis gali būti sumažintas, jeigu reklamose aiškiai nurodomi DI naudojimo etiniai motyvai.

Baek'o et al. (2024) tyrime nustatyta, kad DI turinio autorystės atskleidimo poveikis vartotojų požiūriui į labdaros reklamą ir jų aukojimo ketinimus yra neigiamas – reklamos, kuriose nurodyta, jog jas sukūrė DI, buvo suvokiamos kaip mažiau patikimos. Taip pat, tyrimas parodė, kad tais atvejais, kai respondentai suvokė DI kaip panašų į žmogų, DI autorystės atskleidimo neigiamas poveikis buvo mažesnis, tačiau jei respondentai manė, jog DI yra mažiau panašus į žmogų, tai sustiprino jų neigiamą reakciją. Galiausiai, anot tyrimo rezultatų, DI turinio atskleidimas taip pat turėjo neigiamą poveikį aukojimo sumoms, tačiau šis poveikis buvo mažesnis, kai DI buvo suvokiamas kaip „žmogiškas“.

Bui et al. (2024) nagrinėjo DI sukurtų vizualizacijų suvokiamo autentiškumo įtaka turistų ketinimams lankytis tam tikrose vietose, kai turinio autorystė pažymėta ir nepažymėta. Tyrimas pristatė naują DI autentiškumo (angl. *AI-thenticity*) sąvoką, kuri pabrėžė, kaip turistai suvokia DI sugeneruotą turinį ir jo poveikį jų elgesiui. Rezultatai atskleidė, kad DI sukurtų vizualizacijų suvokiamas autentiškumas reikšmingai didina pasitikėjimą turiniu ir respondentų ketinimus lankytis atitinkamoje vietoje, tačiau DI turinys buvo laikomas mažiau autentišku nei žmogaus kurtas, ypač jei buvo nurodyta, kad nuotrauką sukūrė DI. Taip pat, DI autorystę žyminčios etiketės nebuvimas didino DI sukurtų vizualizacijų suvokiamą autentiškumą.

Chen'o et al. (2025) tyrime buvo analizuojama, kaip DI atskleidimas veikia vartotojų įsitraukimą stebint trumpus vaizdo įrašus apie produktus bei jų naudojimą. Eksperimentas parodė, kad DI atskleidimas turi dvejopą poveikį: viena vertus, vartotojai labiau įsitraukia, kai žino, kad turinį sukūrė DI, nes tai kelia jų smalsumą, kita vertus, turinio autorystės atskleidimas mažina suvokiamą turinio kokybę, nes vartotojai abejoja DI gebėjimu sukurti aukštos kokybės turinį.

Socialiniuose tinkluose vartotojų atsakas į DI sukurtą turinį dažniausiai yra neigiamas, kai atskleidžiama, jog turinį sukūrė DI, lemiantis mažesnę įsitraukimą. Autoriai Brüns'as ir Meißner'is (2024) tyrime analizavo, kaip DI naudojimas kuriant turinį socialiniuose tinkluose veikia vartotojų suvokiamą prekių ženklo autentiškumą bei elgseną. Eksperimento, kuriame buvo lyginamas DI sugeneruotas ir žmogaus sukurtas prekių ženklo turinys, rezultatai parodė, kad DI naudojimas mažina prekių ženklo autentiškumo suvokimą, neigiamai veikdamas įrašų patikimumą, vartotojų ketinimus, susijusius su elektronine „iš lūpų į lūpas“ komunikacija (angl. *electronic word-of-mouth*) ir lojalumą prekių ženklui. Eksperimento dalyviams vertinant DI kurtus ir žmogaus kurtus vaizdinius socialinių tinklų įrašuose, paaiškėjo, kad respondentai dažnai negalėjo atskirti DI kurtų nuo žmonių kurtų vaizdų, o DI autorystės atskleidimas neigiamai veikė prekių ženklo autentiškumą bei dalyvių vertinimus, lyginant su tradiciniu, žmogaus kurtu turiniu. Tiriant DI naudojimą kaip pagalbą žmonėms kuriant turinį, autoriai atrado, jog DI, kaip pagalbos kuriant turinį, naudojimas turėjo mažesnę neigiamą poveikį nei visiškai turinio kūrimo automatizacija, tačiau vis tiek neigiamai veikė vartotojų suvokiamą autentiškumą. Šio tyrimo išvados rodo, kad DI naudojimas gali pakenkti prekių ženklo autentiškumui, ypač kai turinio kūrimo procesas yra visiškai automatizuojamas.

Tyrimo, nagrinėjančių kaip DI sukurtas turinys ir jo žymėjimas veikia vartotojų įsitraukimą socialiniame tinkle „TikTok“, Carney et al. (2024) pasitelkė mišrų metodą, apimančią tiek realių duomenų analizę, tiek eksperimentus. Pirma, buvo išanalizuoti vartotojų elgsenos duomenys socialiniame tinkle „TikTok“ – tyrėjai stebėjo, kaip keitėsi vartotojų įsitraukimo rodikliai (pavyzdžiui, paspaudimai „patinka“) 3 mėnesiai prieš ir 6 mėnesiai po DI turinio žymėjimo politikos įsigaliojimo sąveikaujant su turiniu, kuris buvo pažymėtas kaip sukurtas DI ir su turiniu, kuris nebuvo pažymėtas kaip sukurtas DI. Antra, siekiant patikrinti priežastinius ryšius tarp DI žymėjimo ir vartotojų įsitraukimo, buvo atlikti eksperimentai. Jų metu dalyviams buvo parodomi „TikTok“ platformoje įprastai matomo formato vaizdo įrašai, tačiau vienu vaizdo įrašų apraše (angl. *caption*) buvo matoma DI sukurtu turinio etiketė „sugeneruota DI“ (angl. *AI-generated*), o kituose – ne. Toks eksperimentas įgalino tyrėjus išanalizuoti, kaip DI turinio žymėjimas socialiniuose tinkluose veikia vartotojų elgseną, nepriklausomai nuo kitų galimų veiksnių, tokių kaip turinio kokybės ar vizualinio turinio patrauklumo.

Tyrime suformuluota pagrindinė hipotezė – DI turinio žymėjimas sumažina vartotojų įsitraukimą ne dėl objektyvių turinio kokybės ar autentiškumo suvokimo skirtumų, bet dėl mažesnio emocinio ryšio su turinio kūrėju. Remiantis ankstesniais tyrimais, autoriai nustatė, kad vartotojai dažnai užmezga vienus socialinius ryšius su turinio kūrėjais, kurie laikui bėgant didina jų įsitraukimą į skelbiamą turinį. Kai vartotojai suvokia, kad turinys sukurtas DI, šis socialinis ryšys gali susilpnėti, nes DI laikomas mažiau autentišku ir mažiau atspindinčiu kūrėjo asmenybę ar tikrąsias emocijas. Be pagrindinės hipotezės, tyrėjai taip pat nagrinėjo galimybę, kad skirtingos DI žymėjimo formos galėtų sumažinti arba padidinti vartotojų įsitraukimą. Pavyzdžiui, vienuose eksperimentuose DI žymėjimas buvo vertinamas kaip neutralus, mat naudojo „TikTok“ platformoje realizuotas DI turinio žymėjimo etiketes, tačiau kituose buvo naudojami pozityviai skambantys žymėjimai, pavyzdžiui, „Pažiūrėkite, kokią įspūdingą efektą sukūrė DI!“ (angl. „*Check out this awesome #AI effect!*“). Hipotezė suponavo, kad pozityvus žymėjimas gali padėti neutralizuoti DI etiketės neigiamą poveikį.

Tyrimo rezultatai patvirtino pagrindinę hipotezę: DI žymėjimai reikšmingai sumažino vartotojų įsitraukimą į turinį. „TikTok“ duomenų analizė parodė, kad įrašai, kuriuose DI turinys buvo pažymėtas, sulaukė vidutiniškai 11,9 proc. mažiau „patinka“ paspaudimų nei tie patys įrašai be žymėjimo. Šis skirtumas išliko net atsižvelgus į kitas galimai įtakos turinčias aplinkybes, tokias kaip turinio žiūrėjimo trukmė, aprašyme naudojami raktažodžiai (angl. *hashtags*) ir turinio raiška. Eksperimentiniai tyrimai dar labiau sustiprino šią išvadą: net kai dalyviai matė visiškai identišką turinį, tie, kurių matomas turinys buvo pažymėtas kaip sukurtas DI, nurodė mažesnę ketinimą įsitraukti į tokį turinį. Be to, kai jų buvo paprašyta pasirinkti, kurį įrašą labiau norėtų pažymėti „patinka“, dalyviai žymiai rečiau rinkosi įrašą su DI atskleidimu (tik 35,1 proc. rinkosi DI žymėtą įrašą, palyginti su 64,9 proc. be atskleidimo). Eksperimentai taip pat parodė, kad skirtingos DI turinio žymėjimo formos (pozityvios ar neutralios) neturi reikšmingo poveikio įsitraukimo ketinimams – DI žymėjimas bet koku atveju mažino vartotojų ketinimą įsitraukti, mat vartotojai, matydami DI žymėjimą, jautė mažesnę ryšį su turinio kūrėju.

Atlikti tyrimai atskleidžia, kad DI sukurtų turinio autorystės žymėjimas reikšmingai veikia vartotojų atsaką, tačiau šis poveikis nėra vienodas. Kai kurie tyrimai rodo, kad DI autorystės atskleidimas dažnai mažina vartotojų pasitikėjimą turiniu, prekių ženklo autentiškumo suvokimą (Brüns & Meißner, 2024), ketinimus pirkti (Luo et al., 2019; Powers et al., 2023), keliauti (Bui et al., 2024), stebėti tokį turinį pakartotinai (Chen et al., 2025) ar jį įsitraukti (Carney et al., 2024), mat vartotojai vertina DI turinį kaip mažiau autentišką (Bui et al., 2024) ar kokybišką (Chen et al., 2025). Tiesa,

neigiamas autorystės žymėjimo poveikis gali būti sumažintas, kai aiškiai komunikuojami etiniai DI naudojimo motyvai (Arango et al., 2023; Baek et al., 2024), tačiau tik labdaros komunikacijos atvejais.

Remiantis apžvelgtais tyrimais, vartotojų atsakas į dirbtinio DI sukurto turinio žymėjimą gali būti vertinamas kaip daugialypis reiškiny, apimantis kognityvinius, emocinius, elgsenos ir afektyvinius komponentus. Tolesniuose šio darbo skyriuose bus koncentruojamasi į tris vartotojų atsako į DI turinį pasireiškimo aspektus: šaltinio ir turinio atitikimo suvokimą, patikimumo suvokimą bei įsitraukimo ketinimus. Šie pasireiškimai pasirinkti kaip esminiai aspektai, leidžiantys nuosekliai analizuoti, kaip vartotojai vertina DI sukurtą komunikaciją ir kokius psichologinius bei elgsenos procesus ji sukelia.

2.3.1. Suvokiamas šaltinio ir turinio patikimumas

Kadangi DI autorystės žymėjimas gali paveikti vartotojų pasitikėjimą, vienas svarbiausių veiksnių, padedančių paaiškinti jų atsaką į tokį turinį, yra jo patikimumo suvokimas. Patikimumas (angl. *credibility*) – vienas svarbiausių komunikacijos tyrimų konstrukty, tradiciškai suvokiamas kaip turinio autoriaus ar šaltinio savybė, t. y., patikimumas anksčiau buvo siejamas su tuo, ar autorius yra patikimas (Appelman & Sundar, 2016). Patikimumas buvo suprantamas kaip bendras vertinimas, apimantis šaltinio autoritetą, rinkodaros kanalų reputaciją bei jų tarpusavio sąveiką. Tačiau šiuolaikiniai tyrimai, atsižvelgdami į skaitmeninių kanalų įvairovę, pabrėžia, kad patikimumas nėra vien su turinio autoryste siejama sąvoka. Mokslininkai, tokie kaip Metzger'is et al. (2003), teigia, kad medijos žinutės patikimumas gali būti veikiamas ne tik šaltinio savybių, bet ir kitų veiksnių, tokių kaip žinutės struktūra ar informacijos pateikimo kanalas. Remiantis šia įžvaga, patikimumas imtas skirstyti į tris skirtingas sąvokas:

- šaltinio patikimumą (angl. *source credibility*);
- kanalo patikimumą (angl. *medium credibility*);
- žinutės patikimumą (angl. *message credibility*).

Visgi, prieš aptariant patikimumo sąvokas, svarbu apibrėžti terminą pasitikėjimas, kuris padeda suprasti, kaip vartotojai suvokia ir vertina žiniasklaidą. Komunikacijos tyrimuose sąvokos patikimumas ir pasitikėjimas yra dažnai sugretinamos, mat pasitikėjimo teorinis pagrindas dažnai remiasi būtent medijų patikimumo tyrimų rezultatais (Kohring & Matthes, 2007). Remiantis Hardin'u (2001), pasitikėjimas gali būti apibrėžiamas kaip santykis, kuriame pasitikintysis tikisi, kad kitas asmuo elgsis patikimai, nes yra suinteresuotas išlaikyti santykį ir atliepti pasitikėjusiojo interesus arba dėl moralinio įsipareigojimo. Tai reiškia, jog žmonės pasitiki tikėdamiesi, jog kiti žmonės turi stiprius moralinius principus ir yra pasiryžę laikytis tam tikrų įsipareigojimų, net jei tai jiems nenaudinga.

Šaltinio patikimumas dažniausiai apibrėžiamas, kaip vartotojų pasitikėjimu šaltiniu ir jo gebėjimu pateikti teisingą informaciją tam tikru metu (Johnson et al., 2024). Anot Hovland'o et al. (1953) šaltinio patikimumas yra dvilypė sąvoka, kurią sudaro patikimumas ir kompetencija (angl. *expertise*). Patikimumas apibrėžiamas kaip šaltinio intencija pateikti sąžiningą ir teisingą informaciją, o kompetencija – kaip tam tikros srities gilus išmanymas (Ecker & Antonio, 2021).

Šaltinio patikimumas dažniausiai matuojamas naudojant validuotas skales, tokias kaip McCroskey'io ir Teven'o (1999) sukurta šaltinio patikimumo skalė, kuri vertina patikimumą, kompetenciją ir geranoriškumą (angl. *goodwill*), susijusį su vartotojų suvokiamu šaltinio siekiu rūpintis jų gerove (Johnson et al., 2024). Taip pat dažnai naudojama Ohanian'o (1990) skalė, apimanti patrauklumą

(angl. *attractiveness*), patikimumą ir kompetenciją (Jenkins et al., 2020). Šios skalės leidžia tiksliai įvertinti vartotojų reakcijas į skirtingus šaltinius įvairiuose komunikacijos kontekstuose.

Šaltinio patikimumas veikia vartotojų atsaką į informaciją įvairiais lygmenimis. Pavyzdžiui, Jenkins'o et al. (2020) literatūros apžvalga parodė, kad ekspertų pateikta informacija yra suvokiama kaip patikimesnė nei kitų vartotojų pateikiama informacija. Šis efektas ypač ryškus sveikatos komunikacijoje, kur informacijos autorystė dažnai lemia jos priimtinumą. Kaip rodo tyrimai, tarp įtaką patikimumui darančių turinį skelbiančio autoriaus elementų – ir suvokiamas panašumas su vartotojais (angl. *perceived similarity*), kuris reikšmingai prisideda prie vartotojų pasitikėjimo reklamomis (Lou & Yuan, 2019). Visgi, socialiniuose tinkluose šaltinio patikimumas tampa mažiau aktualus, nes vartotojai dažnai susiduria su didžiuliu informacijos srautu, kuriame ne visada aišku, kas yra žinutės autorius, o daugelis vartotojų neturi pakankamai įgūdžių kritiškai vertinti šaltinio reputaciją ar autoritetą, jei anksčiau su juo nėra susidūrę.

Tuo pačiu, DI sukurto turinio atsiradimas lemia patikimumo ir pasitikėjimo sąvokų kompleksiskumą, mat išnyksta aiškos ribos tarp turinio kūrėjo ir jo platintojo, kadangi šaltinio patikimumo teorijoje pasitikėjimas buvo nukreiptas į konkretų šaltinį – turinį publikuojančią organizaciją ar asmenį – kuris buvo laikomas atsakingu už turinio kokybę, intencijas ir tikslumą. Tačiau DI generuojamo turinio atveju vartotojams dažnai nėra aišku, kas iš tikrųjų sukūrė turinį – žmogus ar DI. Tokiose situacijose kyla klausimas, ar vartotojai pasitiki įrašą paskelbusiu vartotoju, ar technologija, kuri tą turinį sugeneravo. Tai keičia patikimumo vertinimo pagrindus ir reikalauja naujo požiūrio į šaltinio patikimumą, apibrėžiant, kam turėtų būti priskirta atsakomybė už kuriamą informaciją. Dėl to DI technologijų vartojimo atvejais šaltinio patikimumas nebegali būti vertinamas vien tik pagal tradicines šaltinio ypatybes – jis apima ir vartotojų santykį su pačia technologija bei jų suvokimą apie jos intencijas ir veikimo principus.

Pasitikėjimui informacija didelę įtaką daro ne tik pasitikėjimas jos šaltiniu, bet ir informacijos perdavimo priemonės patikimumas (Shamdasani et al., 2001). **Kanalo patikimumo** (angl. *medium credibility*) teorija teigia, kad medijos reputacija ir patikimumas gali stipriai veikti vartotojų pasitikėjimą pateikiama informacija, nepriklausomai nuo pačios žinutės turinio (Choi & Rifon, 2002). Kanalo patikimumą veikia reputacija, kuri nurodo, ar pasirinkta medija yra žinoma kaip prestižinis, patikimas informacijos šaltinis, bei kanalo ir skelbiamos informacijos suderinamumas.

Kanalo patikimumas gali būti matuojamas remiantis tyrimo dalyvių atsakymais apie kanalą pagal tokius kriterijus kaip „patikimas“ ar „nepatikimas“ (angl. *believable, unbelievable* arba *credible, not credible*) ir „įtikinamas“ ar „neįtikinamas“ (angl. *convincing, unconvincing*) (MacKenzie & Lutz, 1989, Choi & Rifon 2002), naudojant septynių balų Likerto skalę, kuriose dalyviai atsako į šešis klausimus, susijusius su interneto naudojimo įpročiais, suvokiamu tinklalapio patikimumu, pasitikėjimu internetu valdančiomis institucijomis, pasitikėjimu kitais interneto naudotojais, interneto naudingumu ir privatumo apsauga (Lucassen & Schraagen, 2012) ar kitais metodais.

Kanalo patikimumas tai pat turi reikšmingą poveikį vartotojų atsakui. Tyrimai parodė, kad aukšto patikimumo kanalai gali sustiprinti žinutės įtikinamumą, skatinti teigiamą vartotojų požiūrį į pateikiamą informaciją ir didinti pasitikėjimą reklamuojamu prekių ženklu. Be to, Choi'us ir Rifon'as (2002) nurodo, kad medijos patikimumas gali paveikti vartotojų pirkimo ketinimus ir informacijos priėmimo procesus.

Vis dėlto, kanalo patikimumas yra neatsiejamas nuo turinio atributų. Anot Shamdasani'o et al. (2001) reklamuojant aukšto įsitraukimo (pirkimo sprendimas reikalauja ilgesnio svarstymo, tikėtina, dėl produkto didelės kainos) produktus, reklaminio vizualo efektyvumas priklauso nuo vizualo turinio ir tinklalapio tematikos sąsajos, o kanalo reputacija daro įtaką tik tada, kai tinklalapio turinys yra aktualus reklamuojamam produktui. Kitaip tariant, jei tinklalapio turinys neatitinka reklamos tematikos, net ir aukštos reputacijos tinklalapis neskatina teigiamo vartotojų atsako. Reklamuojant žemo įsitraukimo produktus (pirkimo sprendimas nereikalauja ilgo svarstymo dėl produkto mažesnės kainos), tinklalapio reputacija tampa lemiamu veiksnio bei gali kompensuoti reklamos turinio ir tinklalapio atitikimo trūkumą. Tačiau, jei tinklalapio reputacija yra žema, reklamos ir tinklalapio suderinamumas tampa svarbiu, mat vartotojai remiasi šiais elementais siekdami įvertinti reklamos patikimumą (Shamdasani et al., 2001). Todėl socialiniuose tinkluose, kur dažnai susiduriama su turiniu, kurio kanalo reputacija gali būti neaiški, svarbesniu veiksnio tampa pačios žinutės patikimumas, nes vartotojai daugiau dėmesio skiria turinio kokybei nei kanalo ar šaltinio autoritetui analizuoti.

Turinio patikimumas (angl. *message credibility*) apibrėžiamas kaip individualus komunikacijos turinio teisingumo vertinimas (Appelman & Sundar, 2016). Šis konceptas atskiriamas nuo šaltinio patikimumo ir kanalo patikimumo, nes dėmesys sutelkiamas į patį žinutės turinį, o ne į tai, kas ją sukūrė arba kokioje platformoje pavišino. Tai reiškia, kad žinutės patikimumas nėra nulemtas šaltinio reputacijos ar kanalo tipo, bet labiau susijęs su turinio informacine verte (Lou & Yuan, 2019), tikslumu, klaidų stoka, teisingumu, autentiškumu ar įtikinamumu (Appelman & Sundar, 2016). Šis konceptas tampa ypač svarbus analizuojant žinutes socialiniuose tinkluose, kur dažnai susiduriama su turiniu, kurio autoriaus ar kanalo reputacija gali būti neaiški arba nežinoma.

Appelman'o ir Sundar'o (2016) teigimu, žinutės patikimumas yra daugiadimensinis bei gali priklausyti nuo žinutės tikslumo (angl. *accuracy*), autentiškumo (angl. *authenticity*) ir įtikinamumo (angl. *believability*). Šios dimensijos leidžia objektyviai įvertinti žinutės patikimumą nepriklausomai nuo autoriaus ar kanalo savybių. Turinio patikimumo matavimas atliekamas naudojant klausimus, kuriuose dalyviai vertina, kiek pateikta žinutė atitinka šias charakteristikas skalėje nuo „labai blogai apibūdina“ iki „labai gerai apibūdina“. Patikima žinutė yra suvokiama kaip labiau įtikinanti, ji didina pasitikėjimą pateikiama informacija, skatina teigiamą vartotojų atsaką, įskaitant dalijimąsi turiniu ir teigiamą įsitraukimą (Shamim & Islam, 2022). Turinio informatyvumas, lemiantis vartotojų pasitikėjimą žinute, gali teigiamai veikti pirkimo ketinimus (Lou & Yuan, 2019; Shamim & Islam, 2022).

Tyrimuose, analizuojančiuose DI naudojimą rinkodaroje, pasitikėjimas tampa ypač svarbus nagrinėjant DI sugeneruoto turinio suvokimą. Nors patikimumas ir pasitikėjimo dažnai vartojami kaip persidengiantys konstruktai, pasitikėjimas apima santykiu grįstą vertinimą, kai vartotojai tiki, kad kiti asmenys veikia remdamiesi moraliniais principais nesiekdami asmeninės naudos. Visgi šis pasitikėjimo mechanizmas DI generuojamo turinio atveju tampa komplikuoatas, mat nebėra aišku, kas yra tikrasis komunikacijos kūrėjas – žmogus ar technologija. Dėl šios priežasties vartotojams gali būti sunku nuspręsti, ar jie turėtų pasitikėti turinį skelbiančiu asmeniu, ar pačia technologija, kuri tą turinį sukūrė. Todėl tam ypač tinka patikimumo kaip daugiadimensio konstrukto samprata, leidžianti išskirti skirtingus pasitikėjimo aspektus ir geriau suprasti, kaip vartotojai vertina DI sukurtą turinį neaiškos autorystės kontekste. Šiame projekte suvokiamas patikimumas aiškinamas kaip apimantis šaltinio, kanalo ir pačios žinutės patikimumą. Remiantis Appelman ir Sundar (2016), Metzger et al. (2003) bei Choi ir Rifon (2002) tyrimais, šios trys dimensijos padeda išsamiai suprasti, kaip

vardotojai vertina informaciją, ypač kai ji yra sugeneruota DI. Šaltinio patikimumas šiame tyrime suvokiamas kaip vartotojas vertina šaltinį remdamasis tuo, kiek jis jam atrodo kompetentingas, empatiškas, palankus ir gebantis užmegzti pasitikėjimu grįstą ryšį. Kanalo patikimumas vertinamas pagal medijos reputaciją ir vartotojų pasitikėjimą platforma, kurioje turinys skelbiamas. Visgi, kanalo patikimumas projekte nėra toliau nagrinėjamas, nes socialiniuose tinkluose daug svarbesniu veiksnio tampa pats turinys ir jo šaltinis. Galiausiai, žinutės patikimumas siejamas su turinio tikslumu, autentiškumu ir įtikinamumu, nepriklausomai nuo to, kas jį sukūrė ar kur jis buvo publikuotas. Šis aspektas tampa itin svarbus vertinant DI sugeneruotą informaciją, kai nėra aiškaus žmogaus autoriaus, kurio kompetencija ar ketinimai galėtų būti vertinami, mat vartotojai gali abejoti jo objektyvumu, tikslumu ar turinio kūrimo motyvais. Tyrimai rodo, jog vartotojai dažnai skeptiškai vertina DI generuojamą turinį, ypač kai jis yra aiškiai žymimas kaip ne žmogaus sukurtas (Powers et al., 2023). DI autorystės atskleidimas gali neigiamai paveikti suvokiamą šaltinio ir turinio patikimumą ir sumažinti vartotojų ketinimus įsitraukti į turinį ar pirkti prekes ar paslaugas, nes DI generuojamas turinys dažnai priimamas kaip mažiau autentiškas ir mažiau patikimas nei žmogaus kurtas (Brüns & Meißner, 2024).

2.3.2. Vartotojų suvokiamas šaltinio ir turinio atitikimas

Kadangi šaltinio ir turinio patikimumo vertinimas DI sukurto turinio atžvilgiu tampa vis sudėtingesnis dėl retai aiškiai nurodomos autorystės, svarbu pažvelgti ir į kitą svarbų veiksnį – tai, kaip vartotojai interpretuoja šaltinio ir turinio tarpusavio dermę. Vartotojų pasitikėjimas socialiniuose tinkluose gali būti veikiamas ne tik to, kas ją publikuoja ar kokia tai informacija, bet ir to, kiek turinys atitinka jų susiformavusius lūkesčius apie šaltinį. Todėl tolesniame skyriuje analizuojama šaltinio ir turinio atitikimo samprata, padedanti suprasti, kaip suvokimo schemas lemia vartotojų reakciją į DI generuojamą turinį.

Schemas atitikimo (angl. *schema congruity*) sąvoka yra plačiai taikoma psichologijos ir rinkodaros moksluose, siekiant paaiškinti, kaip vartotojai apdoroja ir vertina informaciją, remdamiesi ankstesne patirtimi ir jau suformuotomis mentalinėmis struktūromis – schemomis (Mandler, 1982; Meyers-Levy & Tybout, 1989). **Schemas atitikimo teorija** teigia, jog tai atvejais, kai pateikiama informacija atitinka vartotojo lūkesčius, susiformavusius per šias schemas, ji yra lengviau suprantama ir paprastai sukelia teigiamą reakciją, mat schemas atitikimas leidžia vartotojui greičiau interpretuoti informaciją, nepatiriant didelės kognityvinės apkrovos. Pavyzdžiui, jei reklama apie sauskelnes pateikiama naudojant kūdikio, kuris dėvi sauskelnes, nuotrauką, vartotojui ši informacija atrodo logiška ir įtikinama, nes dera su jo lūkesčiais ir suvokimo modeliais apie šią produkto kategoriją (Yoon, 2013). Šis atitikimas gali skatinti teigiamą reakciją, nes vartotojui nereikia papildomai stengtis suprasti reklamoje naudojamos nuotraukos prasmės.

Tuo tarpu **schemas neatitikimas** (angl. *schema incongruity*) atsiranda, kai pateikiama informacija neatitinka esamų vartotojo mentalinių struktūrų (Mandler, 1982; Meyers-Levy & Tybout, 1989). Pavyzdžiui, kai reklamuojant internetinę prekybą akcijomis yra naudojama kūdikio nuotrauka, kuri gali prieštarauti vartotojo lūkesčiams (Yoon, 2013). Nors neatitikimas iš pradžių gali sukelti kognityvinį diskomfortą, vidutinio lygio neatitikimai dažnai skatina vartotojų susidomėjimą ir įtraukia juos į aktyvų mąstymo procesą, kurio metu jie stengiasi išspręsti neatitikimus – vartotojas siekia pažinti informaciją ir suprasti, kaip atitinkami elementai yra susiję. Jei neatitikimas yra sėkmingai išspręstas, jis gali sukelti teigiamą atsaką. Tačiau, jei neatitikimas yra pernelyg ekstremalus

arba nepaaiškinamas, tai gali lemti neigiamą požiūrį į pateiktą informaciją – nusivylimą ar susierzinimą (Meyers-Levy & Tybout, 1989, Peracchio & Tybout 1996, Yoon, 2013).

Rinkodaroje schemas neatitikimas gali būti naudojamas kaip veiksmingas reklamos įrankis, sukeliantis vartotojo smalsumą ar skatinantis dėmesio sutelkimą. Vidutinis schemas neatitikimas laikomas strateginiu elementu, skatinančiu vartotoją įsigilinti į reklamos turinį, kartu didinant jos įsimenamumą (Lange & Dahle'n, 2003). Tačiau svarbią reikšmę šiame procese turi vartotojo motyvacija bei informacijos pateikimo būdas. Asmenys, turintys aukštesnę tolerancijos neapibrėžtumui lygį arba daugiau žinių tam tikroje srityje, yra labiau linkę sėkmingai spręsti neatitikimus ir juos vertinti teigiamai (Meyers-Levy & Tybout, 1989). Mažiau žinių turintys vartotojai labiau linkę reaguoti į emocijas, susijusias su suformuotomis schemomis ir neatitikimais. Taip pat, Yoon'o (2013) tyrime pastebėta, kad kai neatitikimas sukelia smalsumą, intrigos jausmą ar nuostabą, vartotojas yra labiau linkęs dėti didesnes pastangas informacijos analizei. Todėl sėkminga neatitikimo strategija priklauso nuo gebėjimo sukurti tinkamą pusiausvyrą tarp neatitikimo sudėtingumo ir jo sprendimo galimybių, pateikiant užuominas, kurios padeda vartotojui išspręsti neatitikimą.

Schemas neatitikimo teorija reklamose buvo analizuota įvairiais aspektais, siekiant suprasti, kaip vartotojai suvokia skirtingus reklamos komponentus: vizualinius elementus, tekstą, prekių ženklo ar produkto atributus bei reklamos sukūrimo šaltinį ir kaip tai gali paveikti vartotojų suvokimą, emocines reakcijas ir sprendimų priėmimą. Harmon-Kizer'io (2014) tyrime buvo nagrinėjama, kaip schemas atitikimas ir neatitikimas tarp garsenybės įvaizdžio ir prekių ženklo veikia vartotojų reakcijas į reklamas. Rezultatai parodė, kad reklamos, kuriose garsenybės įvaizdis mažiau atitiko prekių ženklo įvaizdį, dažnai sulaukė palankesnių vartotojų vertinimų, ypač reklamos atžvilgiu. Tai rodo, kad vidutinis neatitikimas tarp garsenybės ir prekių ženklo sukuria kognityvinį išitraukimą ir padidina reklamos efektyvumą. Tyrimo dalyvių reakcijos atskleidė ir tai, kad mažiau schemą atitinkančiose reklamose vartotojai daugiau dėmesio skyrė garsenybei ir prekių ženklui, tuo tarpu labiau atitinkančiose reklamose daugiau dėmesio buvo skiriama pačiai reklamai. Tai rodo, kad neatitiktis tarp prekių ženklo įvaizdžio ir garsenybės įvaizdžio skatina vartotojus aktyviau apmąstyti jų sąsają, todėl tokios reklamos gali būti efektyvesnės formuojant vartotojų požiūrį į prekių ženklą (Harmon-Kizer, 2014).

Lee'o ir Kim'o (2022) tyrimas parodė, kad, kai reklamoje pabrėžiami pagrindiniai produkto atributai, įprastos, vartotojų lūkesčius atitinkančios spalvos stiprina teigiamą požiūrį į reklamą. Tuo tarpu, kai reklama pristato antrinius atributus, spalvų neatitikimas gali pritraukti vartotojų dėmesį ir sukelti teigiamą reakciją dėl to, jog šis neatitikimas suvokiamas kaip naujovė. Autoriai pabrėžia, kad antrinių savybių pateikimas kartu su nestandartinėmis spalvomis gali būti efektyvus būdas inovatyviems produktams reklamuoti, tuo tarpu tradicinės spalvos kartu su pagrindinėmis produktų ar paslaugų savybėmis tinka siekiant stiprinti pasitikėjimą matoma informacija (Lee & Kim, 2022).

Lee'o et al. (2025) tyrime buvo analizuojamas DI sukurtų sporto reklamų poveikis vartotojų atsakui, ypatingą dėmesį skiriant reklamos **šaltinio ir turinio atitikimui** (angl. *source-message congruence*), reklamose naudojamų modelių efektyvumui bei informacijos apie DI kilmę pateikimo laikui. Dalyviai žiūrėjo DI generuotas reklamas, kuriose buvo naudojami trijų tipų modeliai: tradiciniai žmonės, virtualūs pseudoportretai ir skaitmeniniai dvyniai – tikroviškos virtualios kopijos, sukurtos remiantis realiais asmenimis. Anot autorių (Lee et al., 2025), neatitikimas tarp reklamos pranešimo turinio ir ją sukūrusio autoriaus turi neigiamą poveikį reklamos vertinimui. Kai DI sukurta reklama

akcentuoja žmogiškus fizinius pasiekimus ir emocijas, vartotojai patiria kognityvinį diskomfortą ir į reklamos patikimumą žvelgia skeptiškai, nes toks vartotojų lūkesčių neatitikimas kelia sunkumus suprantant reklamos prasmę ir kontekstą.

Apibendrinant, schemas atitikimo ir neatitikimo teorija suteikia reikšmingų įžvalgų apie tai, kaip vartotojai apdoroja reklamose pateikiamą informaciją. Schemas atitikimas palengvina informacijos supratimą ir dažnai sukelia teigiamą vartotojo reakciją, nes sumažina kognityvinę apkrovą (Mandler, 1982; Meyers-Levy & Tybout, 1989). Tuo tarpu vidutinio lygio schemas neatitikimas gali paskatinti vartotojų smalsumą, įsitraukimą ir didinti reklamos įsimenamumą (Lange & Dahlén, 2003; Yoon, 2013). Tokie neatitikimai skatina aktyvesnį informacijos apdorojimą, ypač jei vartotojai turi pakankamai žinių ir motyvacijos ją interpretuoti (Meyers-Levy & Tybout, 1989). Tyrimai rodo, kad neatitikimai tarp reklamos elementų, pavyzdžiui, prekių ženklo ir garsenybės įvaizdžių, gali padidinti reklamos efektyvumą, skatindami vartotojus giliau analizuoti turinį (Harmon-Kizer, 2014). Be to, kūrybiniai sprendimai, tokie kaip spalvų ar reklamos šaltinio ir pranešimo neatitikimai, turi potencialo sukelti teigiamą reakciją, jei yra tinkamai subalansuoti (Lee & Kim, 2022).

Apibendrinant galima teigti, kad vartotojų suvokiamas šaltinio ir turinio atitikimas šiame projekte suprantamas kaip vartotojo subjektyviai vertinama dermė tarp informacijos turinio ir ją pateikiančio šaltinio. Šis reiškinys grindžiamas schemas atitikimo teorija, kuri teigia, jog vartotojai informaciją apdoroja remdamiesi iš anksto susiformavusiomis mentalinėmis struktūromis – schemomis, o atitikimas tarp šaltinio ir turinio padeda sumažinti kognityvinę apkrovą, sustiprina įtikinamumą ir skatina teigiamą reakciją. Neatitikimas, ypač kai jis stiprus ar nepaaiškinamas, gali sukelti kognityvinį diskomfortą ir skatinti skeptišką požiūrį į turinį. Naujausi tyrimai rodo, kad šis efektas ypač aktualus DI sugeneruoto turinio atvejais: vartotojai tikisi, kad žmogiškas emocijas ar fizinius pasiekimus akcentuojantys pranešimai bus kuriami žmonių, todėl DI įsitraukimas į šias temas kelia neatitikimą tarp lūkesčių ir realybės, apsunkindamas reklamos prasmės interpretavimą (Lee et al., 2025). Taigi, vartotojų suvokiamas šaltinio ir turinio atitikimas apibrėžiamas kaip lūkesčių, kylančių iš turinio prasmės ir šaltinio pobūdžio, suderinamumas, darantis įtaką pasitikėjimui, informacijos interpretacijai bei vartotojų atsakui.

2.3.3. Vartotojų įsitraukimas socialiniuose tinkluose

Kadangi vartotojų suvokimas apie turinio ar jį skelbiančio šaltinio patikimumą ir tarpusavio atitikimą lemia ne tik jų pasitikėjimą, bet ir elgseninį atsaką, kitas svarbus analizės aspektas – vartotojų įsitraukimas socialiniuose tinkluose. Vartotojų įsitraukimas yra plačiai tyrinėjama sąvoka, tapusi esmine analizuojant prekių ženklų ir vartotojų sąveikas šiuolaikinėje rinkodaroje. Remiantis skirtingais teoriniais požiūriais, ši sąvoka aprėpia kognityvinius, emocinius ir elgesio komponentus, kurie kartu formuoja vartotojo patirtį bei santykį su prekių ženklu (van Doorn et al., 2010; Hollebeek et al., 2014; Kumar et al., 2019). Tokie komponentai išreiškiami per klientų įsitraukimo elgseną (angl. *Customer Engagement Behavior, CEB*), kuri mokslinėje literatūroje apibūdinama kaip „vartotojų elgesio išraiška, nukreipta į prekių ženklą ar organizaciją, tačiau neapsiribojanti vien tik pirkimo veiksmiais“ (van Doorn et al., 2010, p. 253). Brodie'is et al. (2011) pabrėžia, kad klientų įsitraukimo elgsena yra glaudžiai susijusi su santykių rinkodaros paradigma bei paslaugų dominavimo logikos (angl. *service-dominant logic*) teorija, kuri išskiria vartotoją kaip aktyvų vertės bendrakūrėjį.

Socialinių tinklų aplinkoje klientų įsitraukimas įgyja ypatingą reikšmę dėl šių platformų interaktyvumo, momentinės sąveikos ir plataus pasiekiamumo. McCay-Peet ir Quan-Haase (2016)

akcentuoja, kad socialinės medijos suteikia vartotojams galimybę ne tik kurti savo skaitmeninę tapatybę, bet ir sąveikauti su kitais tinklo dalyviais, pasitelkiant turinio kūrimo ir sklaidos mechanizmus. Tokia elgsena atitinka klientų įsitraukimo elgsenos esmę, nes atskleidžia ne tik vartotojo sąveiką su prekių ženklu, bet ir jo indėlį į bendros vertės kūrimą skaitmeninėje erdvėje.

Literatūroje galima rasti įvairių vartotojų įsitraukimo socialiniuose tinkluose apibrėžimų, ir nors kiekvienas tyrėjas į šią sąvoką žvelgia iš skirtingos perspektyvos, visos apibrėžtys iš esmės susijusios su vartotojo aktyvumu skaitmeninėje aplinkoje bei jo santykiu su prekių ženklu ir kitais vartotojais. Pavyzdžiui, Le (2018) ir Onofrei'is et al. (2022) daugiausia dėmesio skiria kiekybiškai pamatuojamai elgsenai – vartotojų veiksniams, tokiems kaip „patinka“ paspaudimai, komentarai ar turinio dalijimasis, kuriais galima įvertinti įsitraukimo lygį. Tuo tarpu Kozinets'as (2014) ir Hallock'as et al. (2019) pabrėžia kokybinius aspektus – vartotojų prasminį įsitraukimą, kūrybinę raišką ir ryšių kūrimą su kitais vartotojais, naudojantis prekių ženklo simbolika ar komunikacinėmis priemonėmis. Kiti šaltiniai (Dolan et al., 2016; Schivinski et al., 2016), savo pateiktose apibrėžtyse akcentuoja tiek vidinę vartotojų motyvaciją, tiek išorinį jų elgesį, susijusį su turinio vartojimu, kūrimu ir sklaida.

Išplėstinis literatūroje aprašytų vartotojų įsitraukimo socialiniuose tinkluose sampratų palyginimas pateikiamas 3 lentelėje. Joje apibendrinti skirtingų autorių požiūriai į vartotojų elgseną skaitmeninėje erdvėje, išskiriant pagrindinius terminus, jų apibrėžimus bei paaiškinimus. Šis palyginimas leidžia nuosekliai įvertinti, kaip įvairūs tyrėjai konceptualizuoja vartotojų įsitraukimą, kokie aspektai akcentuojami analizuojant sąveikas tarp vartotojų ir prekių ženklų socialinių tinklų platformose.

3 lentelė. Vartotojų įsitraukimo socialiniuose tinkluose sampratų palyginimas

Autorius(-iai)	Sąvokos pavadinimas	Apibrėžimas ir paaiškinimas
Kozinets, 2014	Socialinis prekių ženklo įsitraukimas (angl. <i>Social Brand Engagement</i>)	Prasmingas ryšys, kūryba ir komunikacija tarp vieno vartotojo ir vieno ar daugiau kitų vartotojų, naudojant prekių ženklą arba su juo susijusią kalbą, vaizdus ir prasmes.
Dolan et al., 2016	Socialinių medijų įsitraukimo elgsena (angl. <i>Social Media Engagement Behavior</i>)	Vartotojo elgesio išraiškos, susijusios su socialinėmis medijomis, kurios neapsiriboja tik pirkimu ir atsiranda dėl motyvaciją skatinančių veiksnių.
Schivinski et al., 2016	Su prekių ženklu susijusios vartotojo internetinės veiklos (angl. <i>Consumer's Online Brand-Related Activities, COBRA</i>)	Su prekių ženklu susijusių internetinių veiksmų rinkinys, kurių intensyvumas priklauso nuo vartotojo sąveikos su socialine medija ir jo dalyvavimo turinio vartojime, kūrime ir sklaidoje.
Boerman et al., 2017	Elektroninė „iš lūpų į lūpas“ komunikacija (angl. <i>electronic word-of-mouth, eWOM</i>)	Socialinė informacijos apie produktus ar prekių ženklus sklaida tarp vartotojų, turinti teigiamą poveikį nuostatoms apie prekių ženklą, pirkimo ketinimams, elgsenai ir lojalumui.
Le, 2018	Internetinis įsitraukimas (angl. <i>Online Engagement</i>)	Vartotojų internetinė elgsena, matuojama pagal vadinamuosius įsitraukimo matuoklius, apimančius tokius internetinius veiksmus kaip naudotojų skaičius, paspaudimų rodikliai, puslapių peržiūros, turinio pamėgimas ir komentarai, priklausomai nuo platformų.
Sitta et al., 2018	Įsitraukimas socialiniuose tinkluose (angl. <i>Social Media Engagement</i>)	Plati sąvoka, apimanti vartotojo ryšį su reklama, medijomis ir prekių ženklais, kuri vystosi kuriant kūrybiškus sprendimus, į kuriuos būtų galima įsitraukti, kuriuose būtų galima dalyvauti, kuriuos verta platinti, apie kuriuos verta kalbėti.
Hallock et al., 2019	Įsitraukimas socialiniuose tinkluose (angl. <i>Social Media Engagement</i>)	Tai, kas vyksta, kai vartotojas kuria santykius su kitais vartotojais ir prekių ženklais. Tai yra daugiau nei tiesiog mygtuko „patinka“

		paspaudimas, komentarai ar įrašų talpinimas socialiniame tinkle. Tai atskleidžia ilgalaikį santykį tarp vartotojų.
Onofrei et al., 2022	Skaitmeninė įsitraukimo elgsena (angl. <i>Digital Behavioural Engagement</i>)	Vartotojo energijos, pastangų ir laiko lygis, skiriamas prekių ženklui per konkrečią vartotojo ir prekių ženklo sąveiką, t. y. turinio, įterpto į vartotojo socialinius tinklus, pamėgimą, dalijimąsi juo ir komentavimą.

Atsižvelgiant į aptartus vartotojų įsitraukimo socialiniuose tinkluose apibrėžimus, šiame projekte bus remiamasi Onofrei'o et al. (2022) pateikiama skaitmeninės įsitraukimo elgsenos samprata, pagal kurią įsitraukimas socialiniuose tinkluose suprantamas kaip vartotojo skiriama energija, pastangos ir laikas, sąveikaujant su prekių ženklu socialiniuose tinkluose – pavyzdžiui, kai vartotojas spaudžia „patinka“, komentuoja ar dalijasi prekių ženklo turiniu savo socialinių tinklų paskyroje. Šis apibrėžimas pasirinktas todėl, kad jis tiksliausiai atspindi vartotojų elgsenos formas, kurios plačiai realizuojamos daugelyje socialinių tinklų platformų: nuo paprasto reakcijos („patinka“ ar kita emocinė išraiška) paspaudimo iki aktyvesnių veiksmų, tokių kaip komentavimas ar turinio dalijimasis.

Vartotojų įsitraukimas socialiniuose tinkluose pasireiškia įvairiomis formomis, kurios tarpusavyje skiriasi intensyvumu ir įgyvendinimui reikalingų pastangų kiekiu. Siekiant struktūruoti vartotojų elgseną socialinių tinklų platformose, dažnai taikomas su prekių ženklu susijusios vartotojo internetinės veiklos (angl. *Consumer Online Brand-Related Activities, COBRA*) modelis, kuris skiria tris pagrindinius įsitraukimo lygius:

- turinio vartojimas – pasyvus elgesys, apimantis turinio stebėjimą, skaitymą ir peržiūrą,
- prisidėjimas prie turinio sklaidos – vartotojų sąveika su turiniu per komentarus, reakcijas ar dalijimąsi,
- naujo turinio kūrimas – aktyviausia įsitraukimo forma, reikalaujanti didžiausių pastangų ir vartotojų kūrybiškumo, kai vartotojai patys generuoja ir skelbia su prekių ženklu susijusį turinį, pavyzdžiui, apžvalgas, tinklaraščio įrašus ar vaizdo įrašus. turinio vartojimą, prisidėjimą prie jo sklaidos ir naujo turinio kūrimą (Schivinski et al., 2016; Dolan et al., 2019; Trunfio & Rossi, 2021).

Galima teigti, jog šiame projekte naudojama įsitraukimo socialiniuose tinkluose sąvoka, kuri remiasi Onofrei'o et al. (2022) tyrimu, geriausiai atliepia vidutinio įsitraukimo lygio veiksmus, apibrėžtus COBRA modelyje. Šiame lygmenyje vartotojai neapsiriboja vien pasyviu turinio vartojimu, tačiau dar ir aktyviai sąveikauja su prekių ženklo turiniu – spaudžia „patinka“, komentuoja, dalijasi ar kitaip reaguoja į turinį. Tokie veiksmai nereikalauja tiek kūrybinių pastangų, kiek naujo turinio kūrimas, tačiau vis dėlto rodo aktyvesnį vartotojo įsitraukimą, išreiškiantį jo poziciją, emocinį ryšį ir norą dalyvauti prekių ženklo komunikacijoje socialinių tinklų erdvėje.

Šio vidutinio įsitraukimo lygio veiksmus tyrėjai įprastai matuoja pasitelkdami kiekybinius rodiklius, tokius kaip „patinka“ paspaudimai, komentarai, dalijimasis ar papildydami pasyvaus įsitraukimo veiksmams – peržiūrų skaičiumi. Tokie duomenys leidžia objektyviai įvertinti vartotojų sąveikos mastą su prekių ženklo turiniu socialiniuose tinkluose ir dažniausiai yra gaunami naudojantis platformų analitiniais įrankiais, tokiais kaip „Facebook Insights“ (Dolan et al., 2019; Trunfio & Rossi, 2021). Alternatyviai taikomi ir normalizuoti įsitraukimo rodikliai, pavyzdžiui, dalijant bendrą įsitraukimo veiksmų skaičių (komentarai, dalijimasis, „patinka“ paspaudimai) iš bendro pasiektų vartotojų skaičiaus (Rossmann et al., 2016; Velazquez-Solis et al., 2024). Toks matavimo būdas leidžia įvertinti santykinį įsitraukimo lygį nepriklausomai nuo auditorijos dydžio. Be to,

eksperimentiniuose tyrimuose įsitraukimas dažnai vertinamas per ketinimų išraišką – respondentų prašoma įvertinti, kiek tikėtina, kad jie imsisi tam tikrų sąveikos veiksmų, pavyzdžiui, paspaus „patinka“, pakomentuos ar pasidalins turiniu (Boerman et al. 2017; Brüns & Meißner, 2024; Carney et al., 2024).

Įsitraukimas socialinių tinklų platformose, ypač pasireiškiantis per tokius veiksmus ar ketinimą juos atlikti kaip turinio pamėgimas, dalijimasis ar komentavimas, pastaruoju metu yra tapęs dažnu akademinio tyrimų objektu. Moksliniuose tyrimuose analizuojami įvairūs veiksniai, darantys įtaką šiems vartotojų elgsenos modeliams, atsižvelgiant į jų sąveiką su prekių ženklo turiniu bei kitais vartotojais skaitmeninėje aplinkoje. Remiantis Onofrei'o et al.(2022) tyrimu, vartotojų įsitraukimą socialiniuose tinkluose lemia keli esminiai veiksniai, susiję tiek su pačia žinute, tiek su jos siuntėju. Tyrimo rezultatai parodė, kad įsitraukimą reikšmingai veikia šaltinio patikimumas ir panašumas tarp siuntėjo ir gavėjo (angl. *source homophily*). Kitaip tariant, vartotojai yra labiau linkę įsitraukti į turinį, kai jis pateikiamas patikimo asmens ir kai turinio kūrėjas atrodo panašus į juos pačius pagal vertybes, pomėgius ar gyvenimo būdą. Nors tyrimas taip pat analizavo informacinę turinio kokybės dimensiją, būtent šaltinio patikimumas ir panašumas tarp siuntėjo ir gavėjo pasirodė esantys stipriausi elgseninio įsitraukimo prognozės veiksniai, o informacinė turinio kokybė neturėjo reikšmingo poveikio.

Vienas iš vis dažniau analizuojamų veiksnių šioje sąveikoje – turinio kilmės atskleidimas, ypač tais atvejais, kai turinį generuoja DI. Naujausi tyrimai rodo, kad DI autorystės atskleidimas gali reikšmingai paveikti vartotojų elgseną – sumažinti jų įsitraukimo lygį, silpninti emocinį ryšį su turinio kūrėju bei mažinti pasitikėjimą turiniu ir jo autentiškumu (Carney et al., 2024; Brüns & Meißner, 2024; Chen et al., 2025). Kadangi socialiniuose tinkluose vartotojų sąveika dažnai grindžiama pasitikėjimu, identifikacija ir socialine atpažintimi, DI kaip neasmeninio kūrėjo įtraukimas kelia iššūkių formuojant teigiamą įspūdį ir palaikant ilgalaikį ryšį su prekių ženklu. Dėl šios priežasties DI turinio žymėjimas tampa ne tik techniniu, bet ir psichologiniu bei komunikaciniu aspektu, galinčiu lemti vartotojų įsitraukimo intensyvumą skaitmeninėje aplinkoje. Šie procesai ir jų intensyvumo lygiai gali svyruoti nuo pasyvaus turinio stebėjimo iki aktyvaus mygtuko „patinka“ paspaudimo, komentarų rašymo, turinio dalijimosi su kitais socialinio tinklo vartotojais ar net naujo turinio kūrimo.

Apibendrinant, šiame projekte vartotojų įsitraukimas socialiniuose tinkluose suprantamas kaip vartotojo skiriamas dėmesys, pastangos ir laikas sąveikaujant su prekių ženklu socialinių tinklų aplinkoje: paspaudžiant „patinka“, komentuojant ar dalijantis prekių ženklo turiniu savo asmeninėje paskyroje. Šis apibrėžimas suformuotas remiantis Onofrei et al. (2022) skaitmeninės įsitraukimo elgsenos samprata, pagal kurią įsitraukimas pasireiškia tokiais veiksmais kaip turinio pamėgimas, komentavimas ir dalijimasis, atspindinčiais vartotojo santykį su prekių ženklu bei norą įsitraukti į bendrą komunikacinę erdvę. Ši samprata projekte susiejama su COBRA modelio vidutinio įsitraukimo lygiu, kuriam būdingas aiškiai išreikštas elgesys, bet dar nereikalaujantis aukšto kūrybinio įsitraukimo. Tokie veiksmai socialinių tinklų tyrimuose dažniausiai vertinami pasitelkiant kiekybinius matavimo metodus, pagrįstus arba elgesio duomenimis iš platformų analitikos sistemų, arba vartotojų saviraiška apie ketinimą įsitraukti – pavyzdžiui, per klausimus apie tai, kiek tikėtina, kad jie paspaustų „patinka“, pakomentuotų ar pasidalintų turiniu. Todėl vartotojų įsitraukimas laikomas elgsenine, tačiau kiekybiškai įvertinama sąveikos forma, kuri atspindi vartotojo dėmesį, motyvaciją ir santykį su turiniu bei jo šaltiniu, ypač svarbią analizuojant DI generuojamo turinio poveikį socialinėje erdvėje.

2.4. Dirbtinio intelekto sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose konceptualus modelis

Šiame poskyryje, remiantis aptartomis teorijomis, paaiškinami ryšiai tarp tiriamojo darbo kintamųjų, suformuluojamos tyrimo hipotezės ir pateikiamas konceptualus modelis, vaizduojantis numatomus šių kintamųjų ryšius.

DI sukurto rinkodaros turinio autorystės žymėjimo poveikis vartotojų suvokiamam turinio ir jo šaltinio atitikimui. Schemos atitikimo teorija teigia, kad vartotojai informaciją apdoroja remdamiesi išankstiniais suvokimo modeliais – schemomis (Mandler, 1982; Meyers-Levy & Tybout, 1989). Kai informacija atitinka šias schemas, ji yra lengviau suprantama, nesukelia kognityvinės įtampos ir paprastai vertinama palankiai. Tačiau kai informacija neatitinka schemų – pavyzdžiui, kai kyla prieštaravimas tarp pranešimo turinio ir jo šaltinio – gali pasireikšti vadinamasis schemos neatitikimas, kuris sukelia vartotojui kognityvinį diskomfortą ar abejonę dėl matomos informacijos patikimumo (Yoon, 2013; Peracchio & Tybout, 1996).

DI sukurto turinio socialiniuose tinkluose atveju autorystės žymėjimas gali sutrikdyti vartotojų informacijos apdorojimo procesą ir paskatinti išankstinius vertinimus apie turinio patikimumą. Pažymėjus, kad turinį sukūrė DI, kyla neatitikimo tarp vartotojų lūkesčių ir turinio pobūdžio rizika – ypač kai turinys susijęs su žmogaus žiniomis, kompetencija ar atsakomybe grįsta veikla (Lee et al., 2025; Zhang & Gosline, 2023). Neigiama vartotojų reakcija gali būti paskatinta neatitikimo tarp lūkesčio (kad socialinių tinklų įrašą kuria žmogus) ir realybės (kad jį sugeneravo DI). Tokia neatiktis, anot Lee et al. (2025), gali trikdyti žinutės suvokimą, skatinti atsargumą ar net nepasitikėjimą šaltiniu, ypač jei nėra aiškiai perteikta, kodėl DI tinkamas tokio turinio kūrėjas. Todėl, remiantis mokslinės literatūros analize, galima kelti prielaidą, kad DI sukurtas ir aiškiai pažymėtas turinys skatina vartotojus suvokti atitikimą tarp pranešimo turinio ir jo šaltinio:

H1: DI sukurtas ir pažymėtas rinkodaros turinys (palyginti su nepažymėtu) sukels didesnį suvokiamą turinio ir jo šaltinio atitikimą.

Moderuojantis turinio sąsajos su DI poveikis ryšiui tarp DI sukurto rinkodaros turinio autorystės žymėjimo ir vartotojų suvokiamo turinio ir jo šaltinio atitikimo. Literatūroje vis dažniau analizuojama, kaip DI autorystė veikia vartotojų suvokimą apie rinkodaros turinį (Kirkby, Baumgarth & Henseler, 2023; Zhang & Gosline, 2023; Lim & Schmalzle, 2024, Shantatula et al., 2024; Lee et al., 2025). Tyrimai rodo, kad vartotojų požiūris į DI sukurtą rinkodaros turinį yra neigiamas, kai DI autorystė yra atskleidžiama (Zhang & Gosline, 2023; Lim & Schmalzle, 2024, Shantatula et al., 2024; Lee et al., 2025), tačiau nėra aiškus vartotojų požiūris analizuojant turinio ir autorystės sąsają.

Remiantis schemos atitikimo teorija, galima daryti prielaidą, kad neigiamas DI autorystės žymėjimo poveikis gali sumažėti, jei rinkodaros turinys yra susijęs su pačiu DI. Kai reklama sukurta DI ir skirta reklamuoti DI technologiją ar DI grįstą inovaciją, tokio turinio šaltinis ir turinys, tikėtina, labiau atitiks vartotojų išankstinius lūkesčius – vartotojai gali būti labiau linkę tokią reklamą priimti kaip „atitinkančią schemą“ (Mandler, 1982; Meyers-Levy & Tybout, 1989). Toks rinkodaros turinys galėtų būti vertinamas palankiai, nes tiek reklaminio turinio autorius, tiek turinys yra susiję tarpusavyje. Šią prielaidą sustiprina Harmon-Kizer'ės (2021) tyrimo rezultatai, parodantys, kad kai reklamos šaltinis ir reklamuojamas objektas yra susiję, vartotojai reklamą vertina kaip natūralią ir patikimą, o tai lemia teigiamą jų atsaką.

Tuo tarpu, jei DI sukurtas turinys reklamuoja produktą ar paslaugą, neturinčią akivaizdžios sąsajos su DI – pavyzdžiui, maisto produktą – gali kilti schemos neatitikimas. Šiuo atveju vartotojams gali būti sunkiau suvokti, kodėl DI yra tinkamas tokio turinio kūrėjas, o neatitikimas tarp lūkesčio ir realybės gali lemti kognityvinį diskomfortą ar net sumažėjusį pasitikėjimą (Lee et al., 2025; Zhang & Gosline, 2023).

Todėl, remiantis schemos atitikimo teorija, galima prognozuoti, kad turinio tematika – t. y. jo sąsaja su DI – gali veikti autorystės žymėjimo poveikį vartotojų suvokiamam šaltinio ir turinio atitikimui. Remiantis argumentacija, formuluojama hipotezė:

H2: DI sukurtas ir pažymėtas rinkodaros turinys (palyginti su nepažymėtu) sukels didesnį suvokiamą turinio ir jo šaltinio atitikimą kai turinys turi sąsają su DI.

Netiesioginė vartotojų suvokiamo turinio ir jo šaltinio atitikimo įtaka vartotojų įsitraukimui socialiniuose tinkluose per suvokiamą šaltinio ir turinio patikimumą. Vartotojų suvokiamas turinio ir jo šaltinio atitikimas gali reikšmingai paveikti tiek šaltinio, tiek turinio patikimumo vertinimus. Nors patikimumas ilgą laiką buvo siejamas su autoriaus savybėmis, pavyzdžiui, jo informuotumu ar moraliniu sąžiningumu (Appelman & Sundar, 2016; Johnson, Butterfuss & Kendeou, 2024), šiuolaikiniai tyrimai rodo, kad vartotojai vertina patikimumą atskirdami šaltinio ir turinio patikimumą (Metzger et al., 2003). Kai tarp turinio ir jo šaltinio kyla neatitikimas – pavyzdžiui, kai DI sukurtas turinys perteikia žinutes apie temas, su kuriomis DI nėra aiškiai susijęs – tai gali kelti abejonių dėl šaltinio tinkamumo. Vartotojai gali vertinti tokį šaltinį kaip stokojantį kompetencijos ar ketinimų veikti sąžiningai – dviejų kertinių šaltinio patikimumo komponentų (Hovland et al., 1953; Ecker & Antonio, 2021). Atsižvelgiant į aukščiau pateiktus argumentus, tikėtina, kad vartotojai, pastebėję aiškų šaltinio ir turinio suderinamumą, vertins šaltinį kaip labiau patikimą, todėl formuluojama hipotezė:

H3: Vartotojų suvokiamas šaltinio ir turinio atitikimas teigiamai veikia vartotojų suvokiamą šaltinio patikimumą.

Be to, suvokiamas šaltinio ir turinio neatitikimas gali daryti neigiamą poveikį ir turinio patikimumui. Kai šaltinis suvokiamas kaip nesusijęs su turiniu, vartotojai gali manyti, kad ir pats turinys yra nekeliantis pasitikėjimo, nesąžiningas ar stokojantis įtikinamumo (Appelman & Sundar, 2016). Tai ypač aktualu socialiniuose tinkluose, kur šaltinio reputacija dažnai lieka nežinoma, todėl didesnis dėmesys skiriamas pačiam turiniui. Atitinkamai, turinio ir šaltinio atitikimas gali paskatinti teigiamą turinio vertinimą, todėl formuluojama hipotezė:

H4: Vartotojų suvokiamas šaltinio ir turinio atitikimas teigiamai veikia vartotojų suvokiamą turinio patikimumą.

Vartotojų įsitraukimas socialiniuose tinkluose yra daugiadimensinis reiškiny, apimantis tiek emocinius, tiek elgsenos komponentus, kurie formuoja aktyvų santykį tarp vartotojo ir prekių ženklo skaitmeninėje erdvėje (van Doorn et al., 2010; Hollebeek et al., 2014). Empiriniai tyrimai rodo, kad vartotojų elgseninis įsitraukimas – įskaitant „patinka“ paspaudimus, dalijimąsi ir komentavimą – yra tiesiogiai susijęs su pasitikėjimu žinute bei jos autoriumi (Onofrei et al., 2022). Kai šaltinio ir turinio neatitikimas sukelia vartotojui kognityvinį diskomfortą, sumažėja motyvacija sąveikauti su turiniu – šios abejonės gali silpninti vartotojų norą parodyti palaikymą ar sąveikauti su turiniu, nes

įsitraukimas, net ir simbolinis („patinka“ paspaudimas ar komentaras), suponuoja tam tikrą pritarimą turiniui (Boerman et al., 2017; Carney et al., 2024). Keliamo hipotezė:

H5: Vartotojų suvokiamas šaltinio ir turinio atitikimas teigiamai veikia vartotojų įsitraukimo ketinimus socialiniuose tinkluose.

Patikimumas yra viena esminių prielaidų vartotojų elgsenai skaitmeninėje erdvėje – jis lemia ne tik pasitikėjimą turiniu, bet ir įsitraukimo veiksmus (Shamim & Islam, 2022). Kai turinys atrodo patikimas, o jį pateikęs šaltinis suvokiamas kaip sąžiningas ir kompetentingas, vartotojams mažiau kyla abejonių dėl žinutės intencijų, todėl jie yra labiau linkę įsitraukti. Keliamos hipotezės:

H6: Vartotojų suvokiamas šaltinio patikimumas teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose.

H7: Vartotojų suvokiamas turinio patikimumas teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose.

Atsižvelgiant į aukščiau pateiktus argumentus, keliamos netiesioginio ryšio hipotezės. Suvokiamas šaltinio ir turinio atitikimas netiesiogiai veikia vartotojų įsitraukimą per turinio patikimumą ir šaltinio patikimumą. Kuo didesnis suvokiamas šaltinio ir turinio atitikimas, tuo didesnis turinio patikimumas, kuris savo ruožtu teigiamai veikia vartotojų įsitraukimą. Atitinkamai, didesnis suvokiamo šaltinio ir turinio atitikimas lemia aukštesnį šaltinio patikimumo vertinimą, kuris teigiamai veikia vartotojų įsitraukimą. Formuluojamos hipotezės:

H8: Šaltinio patikimumas teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose.

H9: Turinio patikimumas teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose.

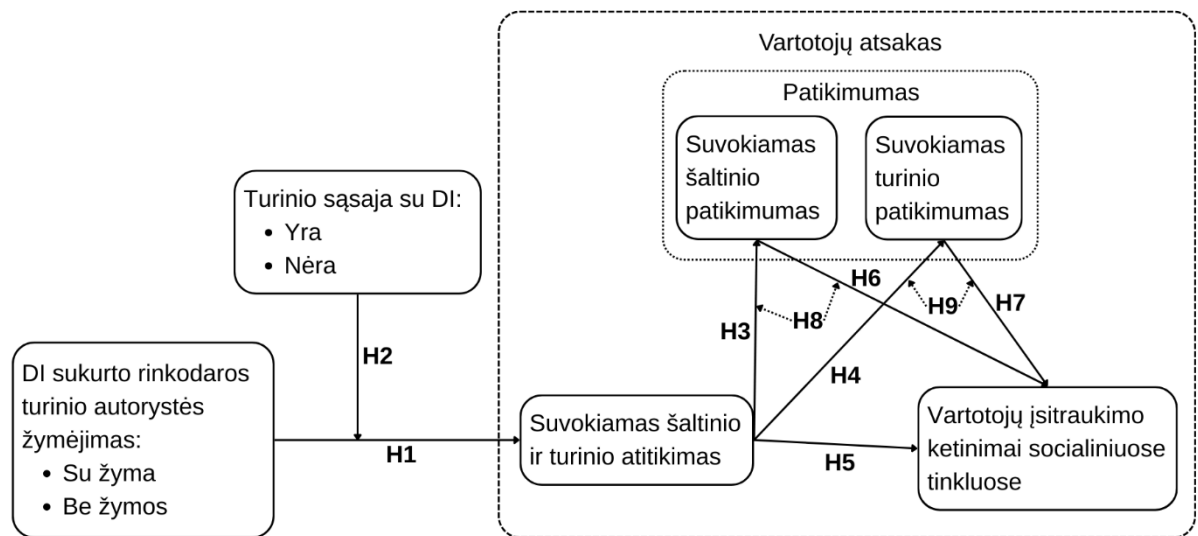
Atlikus išsamią literatūros analizę, teoriškai pagrįsti ryšiai tarp DI sukurto rinkodaros turinio autorystės žymėjimo, turinio sąsajos su DI, šaltinio ir turinio atitikimo bei vartotojų atsako. Pagrindinis dėmesys modelyje skiriamas turinio autorystės žymai – ar turinys yra pažymėtas ar nėra pažymėtas kaip sukurtas DI – kuri gali turėti reikšmingą poveikį vartotojų atsakui.

Modelyje taip pat apibrėžiama, jog suvokiamas šaltinio ir turinio patikimumas daro tiesioginę įtaką vartotojų įsitraukimui socialiniuose tinkluose. Turinys, kuris laikomas patikimu, labiau skatina vartotojų veiksmus, tokius kaip reagavimas į turinį socialinių tinklų platformose realizuotais sprendimais (pavyzdžiui, „patinka“ paspaudimais, turinio dalijimusi ar komentarais). Tuo tarpu turinys, kurio patikimumas kelia abejones, dažniau ignoruojamas arba sulaukia neigiamų reakcijų. Taigi, žymėjimo praktika spėjama turi reikšmingą vaidmenį formuojant tiek vartotojų suvokimą, tiek elgseną.

Modelis pabrėžia galimus moderavimo ir mediavimo ryšius. Turinio sąsaja su DI veikia kaip moderuojantis veiksnys, kuris gali sustiprinti arba susilpninti vartotojų reakcijas, priklausomai nuo to, kaip vartotojai interpretuoja turinio sąsają su DI – ar reklaminio turinio objektas, pavyzdžiui, reklamuojamas produktas ar paslauga, yra susijęs su DI technologijomis. Tuo tarpu šaltinio ir turinio patikimumas veikia kaip priežastinis mechanizmas, tarpininkaujantis tarp suvokiamo šaltinio ir turinio atitikimo bei įsitraukimo ketinimų. Hipotetizuojami ryšiai pabrėžia žymėjimo, turinio sąsajos

su DI bei šaltinio ir turinio atitikimo svarbą ir jų poveikį suvokiamam šaltinio ir turinio patikimumui ir vartotojų įsitraukimui socialiniuose tinkluose.

Remiantis literatūros analize, sudarytas konceptualus modelis (žr. 2 pav.):



2 pav. Konceptualus dirbtinio intelekto sukurto rinkodaros turinio ir autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose modelis

Numatyti konceptualaus modelio ryšiai toliau bus tikrinami empiriškai, siekiant patikrinti hipotetizuojamus ryšius viešųjų mokslo institucijų kuriamų produktų komunikacijoje.

3. Dirbtinio intelekto sukurtos rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose tyrimo metodologija

Šiame skyriuje nagrinėjami empirinio tyrimo metodologiniai aspektai: išskiriamas tyrimo objektas, formuluojami tyrimo tikslas ir uždaviniai, pateikiamas taikytas metodas, aprašomos tyrimui naudotos metodinės bei techninės priemonės, taip pat detalizuojama tyrimo eiga.

Šio tyrimo empirinio tyrimo kontekstas – universiteto mokslinių tyrimų ir jų pagrindu sukurtų produktų bei inovacijų komunikacija socialiniuose tinkluose – pasirinktas neatsitiktinai. Pastaraisiais metais vis daugiau aukštųjų mokyklų ir mokslinių tyrimų institucijų ne tik kuria žinias, bet ir aktyviai diegia inovacijas į rinką per komercializuojamus sprendimus, startuolius ar technologinius produktus (Shah & Pahnke, 2014; Cunningham & Link, 2015; Joshi et al., 2024). Ši tendencija ypač pastebima universitetuose, kurie siekia ne tik akademinio matavimo, bet ir įgyvendinti savo misiją – daryti poveikį visuomenei, ekonomikai ir inovacijų plėtrai. Universitetų mokslinių tyrimų viešinimas ir jų pagrindu sukurtų produktų rinkodara yra esminiai veiksniai, siekiant užtikrinti tyrimų poveikį visuomenei, pritraukti finansavimą ir stiprinti akademinių institucijų reputaciją (Nasirov & Joshi, 2023). Moksliniai tyrimai rodo, kad aukštosios mokyklos, siekdamos išlikti konkurencingos, turi ne tik kurti žinias, bet ir gebėti jas veiksmingai komunikuoti plačiajai visuomenei. Anot Entradas'o (2022), universitetų viešojo komunikacija apie mokslinius tyrimus ir jų rezultatus tampa vis labiau decentralizuota – vis dažniau tyrimų komunikaciją vykdo ne tik universitetų centriniai komunikacijos padaliniai, bet ir patys tyrimų institutai, atskiri tyrimų centrai ar net pavieniai mokslininkai. Tokia decentralizuota komunikacijos struktūra leidžia tyrimų rezultatų sklaidą priartinti prie tikslinių auditorijų, pateikiant informaciją tiksliau, greičiau ir autentiškiau. Be to, aktyvus ir profesionalus tyrimų rezultatų viešinimas, pasitelkiant įvairius komunikacijos kanalus, stiprina universitetų reputaciją, didina pasitikėjimą jų veikla ir skatina tvirtesnių ryšių su pramone, valdžios institucijomis bei kitomis suinteresuotomis grupėmis kūrimą (Entradas, 2022). Anot Schidolski'io et al. (2023), šiandien aukštosios mokyklos privalo taikyti inovatyvias rinkodaros strategijas, kurios atitiktų skaitmeninės komunikacijos realijas ir jaunų auditorijų lūkesčius. Socialiniai tinklai, ypač „Facebook“ ir kitos platformos, tampa pagrindiniais kanalais universiteto žinomumui didinti, reputacijai kurti ir ryšiams su įvairiomis auditorijomis stiprinti.

Patikimumas yra vienas pagrindinių veiksnių, lemiančių vartotojų pasitikėjimą mokslinių tyrimų rezultatais bei jų taikymu praktikoje. Kaip rodo Brewer'io ir Ley'iaus (2013) tyrimas, pasitikėjimas mokslininkais ir kitais mokslo komunikacijos šaltiniais tiesiogiai veikia visuomenės nuostatas apie mokslinius rezultatus. Šią išvadą papildomai patvirtina Brondi'iaus et al. (2021) atliktas tarptautinis tyrimas, parodęs, kad žmonės linkę vertinti mokslo komunikaciją pagal informacijos šaltinio autoritetą ir perduodamų žinių patikimumą. Tyrimo rezultatai atskleidė, kad vartotojai dažniau pasitiki informacija, kuri pateikiama aiškiai, pagrindžiama duomenimis ir skelbiama autoritetingų šaltinių. Naujausi tyrimai (Reif et al., 2024) atskleidžia, kad pasitikėjimas mokslu yra daugialypis reiškinys, susidedantis iš skirtingų dimensijų, tarp kurių – ekspertiskumas (angl. *expertise*) bei geranoriškumas (angl. *benevolence*). Todėl mokslinės komunikacijos patikimumas tampa kritiškai svarbiu veiksnium, siekiant užtikrinti vartotojų pasitikėjimą tiek pačia informacija, tiek jos šaltiniu, o universitetų ir mokslinių tyrimų institucijų gebėjimas kurti patikimą, skaidrią ir auditorijai suprantamą komunikaciją tampa neatsiejama mokslo viešinimo ir praktinio poveikio dalimi.

Aukštojo mokslo institucijose dažnai pasitelkiami pažangūs skaitmeniniai įrankiai, tokie kaip DI (pavyzdžiui, generatyvūs modeliai), kurie naudojami ne tik mokymosi personalizavimui, bet ir

administracinių procesų optimizavimui bei kokybės užtikrinimui. Pastarųjų metų tyrimai rodo, kad aukštosios mokyklos vis aktyviau taiko dirbtinio intelekto technologijas siekdamos pagerinti švietimo kokybę, studentų patirtį ir administracinį efektyvumą (Adel et al., 2025). Kadangi DI technologijos gali būti naudojamos komunikacijos procesų automatizavimui, įskaitant socialinių tinklų turinio kūrimą, nėra aišku, kaip vartotojai suvoktų šaltinio ir turinio patikimumą, ypač jei įrašo autorystė priskiriama DI. Moksliniai tyrimai rodo, kad DI kuriamo turinio atveju svarbi ne tik pati informacija, bet ir jos kilmė bei žymėjimas (Lee et al., 2023; Kirkby et al., 2023). Todėl universiteto kuriamų produktų komunikacija yra ne tik aktualus, bet ir ypač tinkamas kontekstas tirti vartotojų suvokimo aspektus, susijusius su šaltinio ir turinio atitikimu, patikimumu bei įsitraukimo ketinimais, kai komunikacijoje naudojamas DI sugeneruotas turinys. Šis kontekstas leidžia nuosekliai tikrinti tyrimo hipotezes bei prisideda prie aktualaus diskurso apie mokslinės komunikacijos efektyvumą skaitmeninėje erdvėje.

3.1. Tyrimo objektas, tikslas ir uždaviniai

Tyrimo objektas – vartotojų atsakas į dirbtinio intelekto sukurtą rinkodaros turinį apie mokslinių tyrimų pagrindu realizuotus produktus ir turinio autorystės žymėjimą socialiniuose tinkluose.

Tyrimo tikslas – empiriškai pagrįsti DI sukurto rinkodaros turinio apie mokslinių tyrimų pagrindu realizuotus produktus ir jo autorystės žymėjimo poveikį vartotojų atsakui socialiniuose tinkluose.

Uždaviniai:

1. Ištirti, kaip DI sukurtų socialinių tinklų įrašų apie mokslinių tyrimų pagrindu realizuotus produktus žymėjimas veikia vartotojų suvokiamą šaltinio ir turinio atitikimą.
2. Nustatyti, kaip socialinių tinklų įrašų apie mokslinių tyrimų pagrindu realizuotus produktus turinio sąsaja su DI veikia autorystės žymėjimo poveikį vartotojų suvokiamam šaltinio ir turinio atitikimui.
3. Nustatyti vartotojų suvokiamo šaltinio ir turinio atitikimo poveikį suvokiamam šaltinio patikimumui.
4. Nustatyti vartotojų suvokiamo šaltinio ir turinio atitikimo poveikį suvokiamam turinio patikimumui.
5. Nustatyti vartotojų suvokiamo šaltinio ir turinio atitikimo poveikį vartotojų įsitraukimo ketinimams.
6. Įvertinti vartotojų suvokiamo šaltinio ir turinio patikimumo netiesioginį poveikį vartotojų suvokiamam šaltinio ir turinio atitikimui bei įsitraukimo ketinimams socialiniuose tinkluose.
7. Charakterizuoti tyrimo imtį pagal sociodemografinės charakteristikas.

3.2. Empirinio tyrimo metodai

Empirinio tyrimo tipas ir dizainas. Siekiant empiriškai pagrįsti DI valdomų sprendimų sukurto rinkodaros turinio žymėjimo poveikį vartotojų atsakui, pasirinktas kiekybinis tyrimo metodas. Kiekybiniai metodai pasižymi pranašumu prieš kitus tyrimo metodus, nes leidžia tyrėjui objektyviai, nuosekliai ir sistemiškai rinkti bei analizuoti duomenis, tiksliai apibrėžti ir išmatuoti kintamuosius bei sistemiškai pateikti rezultatus (Davis, 2011).

Tyrimo tikslui pasiekti taikomas kiekybinis eksperimento tyrimo dizainas, grįstas scenarijais. Šis metodas yra dažnai naudojamas vartotojų elgsenos tyrimuose siekiant analizuoti priežastinius ryšius tarp kintamųjų (Davis, 2011). Scenarijų taikymas leidžia modeliuoti realias situacijas ir kontroliuoti

tyrimo aplinkybes, siekiant įvertinti, kaip vartotojų atsakas kinta skirtingomis sąlygomis. Eksperimento dizainas panaudojant scenarijus dažnai pasirenkamas ir kituose socialinių mokslų bei rinkodaros tyrimuose, kuriuose nagrinėjamas DI kuriamo ar valdomo turinio poveikis vartotojų suvokimui, įsitraukimui ar elgsenai (Chen et al., 2025; Kirkby et al., 2023; Lee et al., 2025). Ankstesnių mokslinių tyrimų panašiuose kontekstuose taikomi metodologiniai sprendimai pagrindžia eksperimento dizaino tinkamumą nagrinėjamai problemai, leidžia užtikrinti didesnę tyrimo rezultatų patikimumą, ypač hipotetizuojamų priežastinių ryšių atveju.

Tyrime taikomas tarpgrupinis 2x2 eksperimento dizainas remiasi keturiais skirtingais scenarijais, kurie buvo sukurti taip, kad atspindėtų realistiškus socialinių tinklų rinkodaros pranešimus. Kiekvienas dalyvis atsitiktine tvarka priskiriamas vienam iš šių keturių scenarijų, vaizduojamų 4 lentelėje.

4 lentelė. 2x2 tarpgrupinio eksperimento dizaino scenarijų struktūra.

Manipuliacija		DI sukurto rinkodaros turinio autorystės žymėjimas	
		Yra	Nėra
Turinio sąsaja su DI	Yra	Pažymėtas DI turinys, turintis sąsają su DI	Nepažymėtas DI turinys, turintis sąsają su DI
	Nėra	Pažymėtas DI turinys, neturintis sąsajos su DI	Nepažymėtas DI turinys, neturintis sąsajos su DI

Tokiu būdu sukuriamas struktūriškai subalansuotas dizainas, leidžiantis sistemiskai įvertinti, kaip žymėjimo buvimas ar nebuvimas ir turinio kontekstas keičia vartotojų atsaką. Šis metodas leidžia ne tik atskleisti priežastinius ryšius tarp manipuliuojamų kintamųjų ir priklausomųjų kintamųjų, bet ir, remiantis tyrimo rezultatais, generuoti įžvalgas apie DI skaidrumo svarbą rinkodaros komunikacijoje socialiniuose tinkluose.

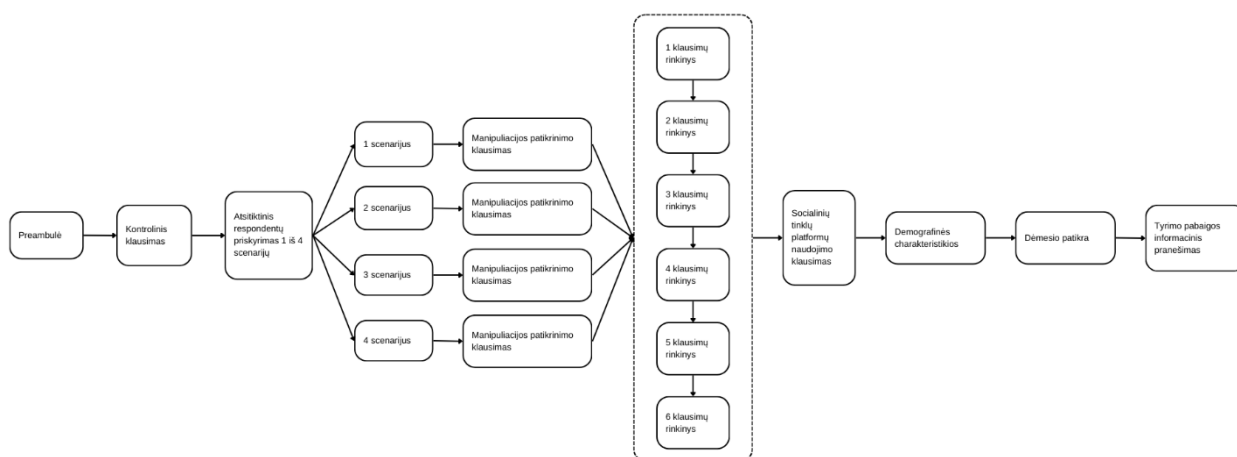
Duomenų rinkimo metodas. Atsižvelgiant į numatytą tyrimo tikslą bei prieinamus išteklius, duomenų rinkimo būdu pasirinkta internetinė anketinė apklausa. Šis metodas sudarė sąlygas pasiekti platų respondentų ratą bei efektyviai surinkti standartizuotus duomenis, tinkamus kiekybinei analizei ir statistiniam apdorojimui.

Šiame tyrime duomenų rinkimui pasirinkta naudoti internetinė apklausų platforma „LimeSurvey“. Šios platformos pasirinkimą lėmė keletas esminių veiksnių, atitinkančių tyrimo dizaino bei metodologinius reikalavimus. Visų pirma, „LimeSurvey“ pasižymi plačiomis galimybėmis kurti sudėtingas apklausas su įvairaus tipo klausimais, šakojimo logika (angl. *branching logic*), atsitiktinai pateikiamais scenarijais ir sąlyginėmis perėjimo taisyklėmis, kas itin svarbu eksperimento dizaino tyrimams, kuriuose siekiama kontroliuoti tyrimo aplinką ir užtikrinti atsitiktinį tyrimo dalyvių skirtingiems scenarijams priskyrimą. Antra, ši atvirojo kodo platforma suteikia galimybę pilnai valdyti tyrimo duomenis – tiek jų kaupimą, tiek saugojimą – atsižvelgiant į etikos ir duomenų apsaugos reikalavimus. Tai leido užtikrinti tyrimo dalyvių konfidencialumą ir laikytis tyrimų etikos gairių. Be to, „LimeSurvey“ palaiko įvairias integracijas su statistinės analizės programine įranga, o gauti duomenys lengvai eksportuojami į SPSS, R ar Excel formatus, kas palengvina tolesnę analizę.

Tyrimo instrumentas. Tyrimo hipotezėms patikrinti buvo pasitelktas konceptualus modelis (žr. 2 pav.), kuriame išskiriamas suvokiamas šaltinio ir turinio atitikimas, suvokiamas šaltinio ir turinio patikimumas bei vartotojų įsitraukimo ketinimai. Modelyje taip pat integruotas nepriklausomas kintamasis – DI sukurto turinio žymėjimas (pažymėtas arba nepažymėtas) bei moderuojantis

kintamasis – turinio sąsaja su DI. Šių kintamųjų pagrindu buvo sudarytas tyrimo instrumentas. Atsižvelgiant į tyrimo tikslus, nuspręsta empirinį tyrimą atlikti Lietuvoje, siekiant pateikti aktualias išvalgas apie nagrinėjamą reiškinį šalies kontekste ir prisidėti prie vietinės mokslinių tyrimų bazės plėtros. Tyrime taikytos matavimo skalės buvo anksčiau naudotos kituose panašaus pobūdžio empiriniuose tyrimuose ir pritaikytos atsižvelgiant į šio tyrimo specifinį kontekstą. Tokia adaptacija leidžia išvengti galimų sunkumų, susijusių su originalių skalių vertimų netikslumais, ir taip sumažina vertimo kokybės įtaką tyrimo rezultatų patikimumui.

Internetinės apklausos instrumentas (žr. 4 priedas) buvo sudaryta iš 15 klausimų blokų (žr. 3 pav.), apimančių eksperimento scenarijus su atitinkama manipuliacija (nepriklausomi kintamieji), dėmesio patikros klausimus, pagrindinius analizuojamus konstruktus bei demografinę informaciją.



3 pav. Tyrimo instrumento schema

Pirmiausia pateikiamas atrankos klausimas, kuriuo tikrinama, ar respondentas yra socialinių tinklų vartotojas – tik tokiu atveju galima tęsti dalyvavimą tyrime. Kiekvienam respondentui atsitiktine tvarka buvo pateikiamas vienas iš keturių scenarijų, kurie skyrėsi pagal du eksperimento dizaino parametrus: turinio žymėjimą kaip DI sukurtą (pažymėtas ar nepažymėtas) ir turinio sąsają su DI tematika (yra ar nėra). Po scenarijų pateikiami manipuliacijos patikros klausimai, kuriais tikrinama, ar dalyvis atpažino, kad turinys sukurtas DI pagalba. Taip pat įtrauktas vienas dėmesio patikros klausimas, padedantis nustatyti, ar dalyvis tinkamai įsigilino į pateiktą scenarijų. Apklausos vidurinę dalį sudaro klausimų blokai, skirti keturiems pagrindiniams konstrukts matuoti skalėmis:

- Šaltinio ir turinio atitikimui,
- Šaltinio ir turinio patikimumui,
- Įsitraukimo ketinimams.

Visi šie konstruktai matuojami Likerto tipo ir semantinio diferencialo skalėmis, o jų formuluotės adaptuotos remiantis anksčiau validuotomis skalėmis. Apklausos pabaigoje pateikiami demografiniai klausimai (amžius, lytis, išsilavinimas, pajamos) ir socialinių tinklų naudojimo įpročiai, padedantys geriau apibūdinti tyrimo imtį bei atlikti papildomą rezultatų analizę pagal sociodemografinius aspektus.

Tyrimo pradžioje respondentams buvo pateikta preambulė – išsami informacija apie tyrimo tikslą, dalyvavimo sąlygas ir asmens duomenų tvarkymo principus, laikantis socialinių mokslų tyrimų etikos reikalavimų. Dalyviams aiškiai nurodyta, jog dalyvavimas tyrime yra visiškai savanoriškas, o jų

atsakymai bus anonimiški ir naudojami tik akademiniams – magistrantūros baigiamojo darbo rengimo – tikslais. Taip pat pažymėta, kad tyrimo duomenys saugomi konfidencialiai, o jų analizė atliekama apibendrinamai, be galimybės identifikuoti konkrečius respondentus. Apskaičiuota anketos pildymo trukmė buvo nurodyta kaip 6–9 minutės.

Siekiant užtikrinti tyrimo duomenų aktualumą ir atitikti tyrimo tikslui, apklausos dalyviams buvo iškelta išankstinė atrankos sąlyga – jie privalėjo būti aktyvūs socialinių tinklų vartotojai. Ši sąlyga buvo tikrinama anketoje pateikiant pirmąjį filtruojantį klausimą („Ar naudojate socialinius tinklus?“), kurio atsakymas buvo būtina sąlyga tolimesniam dalyvavimui tyrime. Respondentams, kurie į šį klausimą atsakė neigiamai, automatiškai buvo pateikiamas paaiškinimas apie neatitikimą tyrimo kriterijams bei prašymas nebetęsti apklausos. Toks sprendimas leido išvengti nereikšmingų atsakymų ir užtikrinti, kad tyrime dalyvautų tik tikslinę grupę atitinkantys asmenys.

„Lime Survey“ programos pagalba, respondentai atsitiktiniu būdu priskiriami vienam iš keturių scenarijų grupių. Šie scenarijai sukonstruoti remiantis 2×2 tarpgrupinio eksperimento dizaino logika, manipuliuojant dviem nepriklausomais kintamaisiais: ar socialinių tinklų įrašas buvo pažymėtas kaip sukurtas dirbtiniu intelektu ir ar įrašas turėjo aiškią teminę sąsają su DI. Kiekvienam respondentui buvo pateiktas tik vienas socialinio tinklo „Facebook“ stiliaus įrašas su paveikslėliu ir tekstu, kurio turinys atitiko vieną iš šių keturių sąlygų. Prieš pateikiant tolesnius klausimus, respondentai buvo raginami įdėmiai įsiskaityti į pateiktą įrašą ir atkreipti dėmesį į jo struktūrą bei informaciją. Tyrimo metu manipuliuojami aspektai nebuvo įvardyti tiesiogiai, siekiant užtikrinti, kad dalyviai į scenarijų reaguotų autentiškai, be išankstinių nuostatų.

Atviru klausimu, pateiktu iškart po scenarijaus, buvo prašoma įvertinti, kiek, jų nuomone, tikėtina, kad pateiktas socialinių tinklų įrašas galėjo būti sukurtas DI. Šis papildomas klausimas buvo įtrauktas siekiant giliau suprasti respondentų mąstymo procesus ir vertinimo kriterijus, kuriais jie rėmėsi formuodami nuomonę apie turinio kilmę. Vertinimui buvo pasitelkta septynių balų Likerto skalė, kuri leido įvertinti respondentų nuomonę. Po šio teiginio pateiktas atviras klausimas, kuriuo buvo prašoma pagrįsti savo atsakymą. Atvirasis formatas leido giliau analizuoti, kokius ženklus ar turinio elementus respondentai laikė svarbiais sprendžiant apie DI įsitraukimą, ir kartu sustiprino tyrimo dalyvių įsijautimą į scenarijų ir manipuliacijos patikimumą.

Pirmasis klausimų blokas buvo skirtas įvertinti, kaip dalyviai suprato ir vertino įrašo bei jo autoriaus suderinamumą – šaltinio ir turinio atitikimą. Buvo pateikti trys teiginiai, padedantys nustatyti, ar turinys buvo laikomas autentišku, ar jame įžvelgta priešara tarp įrašo turinio ir jo šaltinio. Šie klausimai padėjo įvertinti komunikacinės žinutės nuoseklumą bei jos tikslumo suvokimą. Tyrime naudoti kelios konstrukto matavimo skalės, pritaikytos remiantis ankstesniais tyrimais, siekiant užtikrinti matavimo patikimumą ir validumą bei gauti pagrįstus duomenis hipotezėms patvirtinti ar paneigti. Tokiu būdu sudaromos prielaidos tikslesniam rezultatų interpretavimui ir pagrįstų rekomendacijų formulavimui. **Šaltinio ir turinio atitikimas** buvo matuojamas remiantis Lee'o et al. (2025) bei Choi'iaus (2024) metodika, o vertinimui buvo taikyta septynių balų Likerto skalė (1 – „visiškai nesutinku“, 7 – „visiškai sutinku“). Šiame tyrime buvo naudojami trys teiginiai, adaptuoti iš skirtingų autorių darbų:

- „Manau, kad socialinių tinklų įrašas ir jo autorius atitinka vienas kitą“
- „Manau, kad socialinių tinklų įrašo autorius veiksmingai perteikė numatytą žinutę“
- „Manau, kad šis socialinių tinklų įrašas yra susijęs su jo autoriumi“.

Antruoju klausimų bloku siekta įvertinti socialinių tinklų įrašo turinio patikimumą. Dalyvių buvo prašoma įvertinti, kiek įrašas pasirodė patikimas, sąžiningas ir įtikinamas. Šie teiginiai buvo vertinami taip pat septynių balų Likerto skalėje. **Suvokiamas įrašo patikimumas** buvo matuojamas trimis teiginiais, adaptuotais pagal Brüns'o ir Meißner'io (2024) skalę:

- „Šis socialinių tinklų įrašas kelia pasitikėjimą“
- „Šis socialinių tinklų įrašas atrodo sąžiningas“
- „Šis socialinių tinklų įrašas atrodo įtikinamas“

Trečiajame bloke buvo vertinamas **suvokiamas šaltinio patikimumas** naudojant McCroskey'io ir Teven'io (1999) sukurtą semantinio diferencialo skalę, kuri apima tris šaltinio dimensijas: kompetenciją, palankumą ir pasitikėjimą. Dalyviams buvo pateiktos priešingų būdvardžių poros (pavyzdžiui, „Protingas – Neprotingas“, „Sąžiningas – Nesąžiningas“, „Rūpinasi manimi – Nesirūpina manimi“), tarp kurių jie žymėjo savo vertinimą septynių balų skalėje. Instrukcijoje buvo nurodyta pasirinkti skaičių, kuris, jų nuomone, geriausiai atspindi įrašo autorių.

Ketvirtasis klausimų blokas buvo skirtas **įsitraukimo ketinimams** matuoti pagal Onofrei'io et al. (2022) pasiūlytą elgsenos skalę, kurią sudarė trys teiginiai:

- „Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, paspausčiau „patinka“ po šiuo įrašu“
- „Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, palikčiau komentarą po šiuo įrašu“
- „Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, pasidalinčiau juo su kitais“

Atsakymai buvo vertinami septynių balų Likerto skalėje. Šis blokas padėjo įvertinti elgsenos ketinimus, kurie dažnai laikomi svarbiausiu socialinių tinklų rinkodaros komunikacijos efektyvumo rodikliu. Visi tiriamieji klausimų blokai, apimantys įrašo ir šaltinio patikimumo, tarpusavio atitikimo bei įsitraukimo ketinimų vertinimą, buvo pateikiami atsitiktine tvarka, siekiant sumažinti galimą atsakymų šališkumą dėl klausimų sekos efekto.

Tyrimo pabaigoje buvo pateikti klausimai apie socialinių tinklų naudojimo įpročius bei demografiniai klausimai. Respondentų buvo prašoma nurodyti, kiek laiko kasdien skiria socialiniams tinklams, kurias platformas naudoja dažniausiai, taip pat pateikti informaciją apie savo amžių, lytį, išsilavinimą bei pajamas. Siekiant užtikrinti tyrimo rezultatų tinkamumą Lietuvos kultūriniam ir švietimo kontekstui, demografinio klausimo, susijusio su išsilavinimo lygiais, atsakymai tyrimo anketoje buvo suderinti su šalyje egzistuojančiomis mokymosi pakopomis – parengti remiantis Lietuvos Respublikos 2021 m. gyventojų ir būstų surašymo duomenyse pateikiamu išsilavinimo klasifikavimu, skelbiamu Oficialiame Lietuvos statistikos portale (Oficialios statistikos portalas, 2021). Šie duomenys buvo svarbūs siekiant geriau apibūdinti tiriamąją imtį ir įvertinti galimus demografinius skirtumus nagrinėjamuose reiškiniuose.

Dėmesio patikrinimui (angl. *attention check*) buvo pateiktas klausimas, susijęs su pateiktu socialinių tinklų įrašu: „Nurodykite, kokia inovacija buvo pristatoma Jums pateiktame socialinių tinklų įraše“. Dalyviams reikėjo pasirinkti vieną iš pateiktų atsakymo variantų: maisto produktas, DI sistema depresijai atpažinti, biuro kėdė. Šis klausimas padėjo užtikrinti, kad respondentai buvo įdėmūs ir skaitė pateiktą informaciją. Manipuliacijos patikrinimui (angl. *manipulation check*) buvo pasitelktas klausimas, kuriuo siekta įvertinti, ar respondentai pastebėjo eksperimento metu pateiktą ženklumą, susijusį su turinio kilme. Dalyviams buvo pateiktas teiginys: „Įrašas buvo pažymėtas dirbtinio intelekto autorystę nurodančia etikete „AI info“ (liet. dirbtinio intelekto informacija)“, o jų atsakymo variantai buvo: „Taip“, „Ne“ ir „Nežinau“. Šis klausimas leido įvertinti, ar eksperimentinė

manipuliacija – DI autorystės žymėjimo buvimas ar nebuvimas – buvo pastebėta ir suvokta, kaip numatyta tyrimo sąlygomis. Tokiu būdu buvo patikrintas manipuliacijos efektyvumas bei užtikrinta vidinis tyrimo validumas. Išsamus naudotų matavimo skalių ir jų adaptacijų palyginimas pateikiamas 4 priede.

Tyrimo metu naudotų validuotų skalių suvienodinimas į septynių balų Likerto skalę buvo atliktas siekiant užtikrinti duomenų palyginamumą ir analizės nuoseklumą. Nors pradinės skalės buvo validuotos atskiruose tyrimuose ir galėjo skirtis, jų transformavimas į vienodą matavimo formatą leidžia sumažinti matavimo skirtumus tarp kintamųjų, kuriuos lyginama ar jungia bendros analizės. Septynių balų skalė buvo pasirinkta ne tik dėl galimybės tiksliau atspindėti nuomonių niuansus nei naudojant trumpesnes skales (pavyzdžiui, penkių balų), bet ir dėl jos privalumų tyrimo dalyvių požiūriu. Vienodo ilgio ir nuosekliai naudojamos skalės mažina kognityvinę apkrovą, padeda dalyviams greičiau suprasti klausimus (Preston, & Colman, 2000). Be to, nuoseklumas atsakymo formatuose mažina atsitiktinį atsakymų kintamumą, padeda užtikrinti didesnį duomenų patikimumą ir vidinį suderinamumą (angl. *internal consistency*), ypač kai skirtingi konstruktai įvertinami toje pačioje analizėje.

Anketos pabaigoje tyrimo dalyviams buvo pateiktas paaiškinamasis tekstas (angl. *debriefing*), kuriame atskleistas tikrasis tyrimo tikslas, eksperimento manipuliacijos pobūdis bei priežastys, dėl kurių dalis informacijos nebuvo atskleista tyrimo pradžioje. Paaiškinime buvo informuojama, kad visos tyrime naudotos socialinių tinklų žinutės buvo sukurtos DI, tačiau eksperimentinėse sąlygose jos buvo pateikiamos su arba be žymos. Ši manipuliacija atlikta siekiant išmatuoti DI žymėjimo poveikį vartotojų suvokimui ir elgsenos kėtimams. Paaiškinimas buvo parengtas laikantis mokslinių tyrimų etikos ir informuoto sutikimo principų, siekiant užtikrinti tyrimo skaidrumą ir dalyvių teisę žinoti apie tyrimo esmę po jų dalyvavimo. Be to, paaiškinimo pabaigoje nurodyti tyrimo autorės kontaktiniai duomenys, suteikiant galimybę dalyviams susisiekti papildomos informacijos ar tyrimo rezultatų pageidavimo atveju.

Naudotų matavimo skalių ir jų adaptacijų palyginimas tyrime pateikiamas lentelės, kurioje sistemingai palyginami tyrime naudoti pagrindinių konstrukto teiginiai su originaliomis autorių skalėmis anglų kalba, formatu (žr. 3 priedas). Lentelėje nurodomi tiek teiginių formulavimo adaptacijos, tiek taikytos skalės tipas bei naudojimo pagrindimas, remiantis atitinkamais šaltiniais.

3.3. Empirinio tyrimo eigos ir duomenų analizės procedūros

Tyrimo imties dydžio nustatymas. Remiantis *a priori* imties dydžio skaičiavimu, atliktu naudojant G*Power 3.1.9.7 programinę įrangą (Faul et al., 2009), planuojamam tyrimui, kuriame numatytas dviejų veiksnių dispersinės analizės (two-way ANOVA) taikymas, buvo nustatyta, kad esant 0,05 pirmos rūšies paklaidos tikimybei (α), 0,25 efektui (f) bei 0,95 statistinei galiai ($1-\beta$), reikalinga minimali imtis yra 279 tyrimo dalyviai. Atsižvelgiant į galimą dalies atsakymų eliminavimo poreikį (dėl manipuliacijos ar dėmesio patikros neišlaikymo), numatyta, kad faktinė tiriamoji imtis turėtų viršyti minimalų reikalaujamą dydį. Todėl siekta surinkti apie 400 respondentų, kad būtų užtikrintas pakankamas imties dydis ir išlaikytas planuojamo tyrimo statistinis pagrįstumas. Šiame tyrime taikyta netikimybinė, patogioji imtis. Toks pasirinkimas buvo nulemtas ribotų tyrimo autorės finansinių bei laiko išteklių.

Tyrimo imtis. Tyrimo metu buvo surinkti 480 respondentų atsakymai. Respondentai, kurie pateikė neteisingus atsakymus į bent vieną iš dėmesio patikros klausimų, buvo pašalinti iš tolesnės analizės.

Dėmesio patikros klausimai apėmė dvi pagrindines sritis: (1) ar respondentai teisingai identifikavo, apie koki mokslinių tyrimų pagrindu realizuotą produktą buvo kalbama įrašė, ir (2) ar jie tinkamai pastebėjo, ar įrašė buvo nurodyta DI autorystės etiketė. Iš viso dėl neatitikimo dėmesio patikros kriterijams buvo atmesti 108 respondentai, todėl galutinė duomenų analizėje dalyvavusių asmenų imtis sudarė 372 respondentus. Respondentai buvo atsitiktinai priskirti vienam iš keturių eksperimentinių scenarijų, o po duomenų filtravimo grupėse susidarė tokie pasiskirstymai: 1 scenarijus – 96 respondentai, 2 scenarijus – 95 respondentai, 3 scenarijus – 91 respondentas, 4 scenarijus – 90 respondentų.

Tyrimo etika. Tyrimo etikai užtikrinti buvo laikomasi anonimiškumo ir konfidencialumo principų – respondentų pateikti atsakymai nebuvo siejami su jų asmenine tapatybe, o surinkti duomenys naudoti tik magistrantūros baigiamojo projekto rengimo tikslais. Respondentams buvo pateikta išsami informacija apie tyrimą, įskaitant tyrimo tikslą, eigą, numatomą trukmę ir klausimų pobūdį. Kiekvienam respondentui buvo suteikta galimybė bet kuriuo metu nutraukti savo dalyvavimą tyrime be jokio paaiškinimo ar neigiamų pasekmių. Kadangi tyrimo metu respondentams galėjo būti pateikta netiksli informacija apie kai kuriuos eksperimento metu naudojamus elementus (DI sugeneruotas, tačiau nepažymėtas tekstas), tokia praktika buvo iš anksto suplanuota siekiant išvengti atsakymų šališkumo. Anketos pabaigoje dalyviai buvo išsamiai informuoti apie šį apklausos elementą.

Tyrimo instrumento ir manipuliacijų testavimas. Siekiant išvengti galimų metodologinių klaidų, anketos klausimų neaiškumų ar šališkumo, šiame tyrime nuspręsta papildomai pasitelkti didįjį kalbos modelį „ChatGPT“, kaip pagalbinę priemonę tyrimo instrumento peržiūrai ir preliminariam vertinimui. Šis metodas remiasi naujausiomis mokslo įžvalgomis apie didelių kalbos modelių potencialą vartotojų tyrimuose (Sarstedt et al., 2024). Dideli kalbos modeliai gali būti naudingi ankstyvuose tyrimų etapuose, kai atliekamas pirminis klausimyno turinio vertinimas. Remiantis Sarstedt'o et al. (2024), tokie modeliai gali padėti nustatyti dviprasmiškus klausimus, įvertinti jų aiškumą ir net numatyti galimas interpretacijos problemas tarp skirtingų respondentų grupių. Dideli kalbos modeliai taip pat gali būti naudojami įvertinant klausimų kultūrinį neutralumą bei potencialią šališkumo riziką. Kadangi šis tyrimas siekia įvertinti DI sukurto rinkodaros turinio žymėjimo poveikį vartotojų suvokiamam turinio patikimumui ir įsitraukimui, itin svarbu užtikrinti, kad klausimyno formuluotės nesukeltų atsakymų tendencingumo ar klaidingų interpretacijų. Nors dideli kalbos modeliai nėra realių respondentų atitikmuo, jų gebėjimas analizuoti kalbinį turinį ir siūlyti įžvalgas apie potencialius klausimyno trūkumus leidžia tyrėjui dar prieš įvykdant apklausą atlikti preliminarious instrumentų testavimus ir sumažinti riziką, susijusią su metodologiniais netikslumais.

Todėl tyrimo metu buvo pasitelktas DI pagrindu veikiantis kalbos modelis „ChatGPT“, siekiant įvertinti parengtos apklausos anketos kalbinį aiškumą, struktūros nuoseklumą bei tyrimo instrumento atitiktį moksliniams ir etiniams reikalavimams. Tokia metodologinė prieiga leido įvertinti, kaip aiškiai ir logiškai suformuluoti klausimai, ar jie tinkami skirtingo išsilavinimo ir patirties respondentams, ir ar nekyla dviprasmių interpretacijų. Atlikus pirminę analizę, buvo pasiūlyta papildyti kiekvieną scenarijų aiškinamuoju sakiniu, nurodančiu, jog respondentui pateikiamas socialinių tinklų įrašas, į kurį reikia atidžiai įsigilinti prieš atsakant į klausimus. Tai padėjo užtikrinti, kad respondentas suprastų scenarijaus paskirtį ir tinkamai pasiruoštų vertinimui. Kiti pasiūlyti pakeitimai, tokie kaip alternatyvios tyrime naudojamų teiginių formuluotės ar klausimų pertvarkymai, nebuvo įgyvendinti. Jų buvo atsisakyta siekiant išlaikyti tyrimo instrumento suderinamumą su ankstesniais analogiškais tyrimais, kuriuose buvo naudotos standartizuotos teiginių formuluotės. Tokiu būdu buvo užtikrintas rezultatų palyginamumas ir metodologinis nuoseklumas.

Dideli kalbos modeliai gali būti pasitelkti kaip priemonė preliminariam eksperimentinių scenarijų testavimui. Tyrimai rodo, kad tokie modeliai gali padėti įvertinti scenarijų struktūros logiką, vartotojų elgsenos modelių imitavimą ir netgi numatyti, kaip skirtingos formuluotės gali paveikti respondento interpretaciją (Sarstedt et al., 2024). Todėl šiame tyrime taip pat buvo pasitelktas „ChatGPT“, siekiant preliminariai įvertinti eksperimento instrumento struktūrą, scenarijų logiką ir prognozuojamą jų poveikį respondentų elgsenai. Generuodamas sintetinius atsakymus pagal kiekvieną eksperimentinį scenarijų, „ChatGPT“ leido įvertinti, ar manipuliacinės sąlygos sukelia skirtingus atsakymų modelius, ar matavimo priemonės sugeba atskleisti numatytus skirtumus, ir ar priklausomi kintamieji yra jautrūs eksperimentinei manipuliacijai. Tyrimo scenarijai ir manipuliacijos buvo konstruojami remiantis kitų autorių empiriniais pavyzdžiais ir adaptuojami šio tyrimo kontekstui (Yoo et al., 2024), siekiant užtikrinti jų aktualumą ir teorinį pagrįstumą. Pagal sukonstruotą užklausą (angl. *prompt*), buvo sugeneruota 385 respondentų sintetinė imtis ir išanalizuoti gauti duomenys. „ChatGPT“ sugeneruoti duomenys buvo statistiškai analizuoti taikant tiek aprašomąją statistiką (vidurkius, standartinius nuokrypius), tiek inferencinius metodus – vienfaktorinę dispersijos analizę (ANOVA) bei nepriklausomų imčių t-testus. Ši analizė atskleidė, kad tiek DI žymėjimas, tiek turinio sąsaja su DI turėjo reikšmingą poveikį visiems pagrindiniams priklausomiems kintamiesiems: šaltinio ir turinio atitikimui ($p < 0,0001$ ir $p < 0,0001$), šaltinio ir turinio patikimumui ($p < 0,0001$ ir $p < 0,0001$), bei vartotojų įsitraukimo elgsenai ($p < 0,0001$ ir $p < 0,0001$). Siekiant įsitikinti, kad manipuliacija dėl turinio sąsajos su dirbtiniu intelektu buvo efektyvi, į tyrimo klausimyną buvo įtrauktas papildomas klausimas: „Jūsų nuomone, kiek šiame įrašė pristatoma inovacija yra susijusi su dirbtiniu intelektu?“. Sintetinių duomenų analizė parodė, kad vartotojai, kurie matė įrašus su aiškia turinio sąsaja su DI (grupės 1 ir 2), įvertino šią sąsają kaip stipresnę nei tie, kurie matė įrašus be aiškos sąsajos su DI (grupės 3 ir 4). Grupėse, kurios matė įrašus su turinio sąsaja su DI, vidurkiai buvo 4.14 (grupė 1, SD = 1.90) ir 4.22 (grupė 2, SD = 2.13), tuo tarpu grupės, kurios matė įrašus be sąsajos su DI, vidurkiai buvo 3.80 (grupė 3, SD = 2.04) ir 3.58 (grupė 4, SD = 1.98). Atliekant nepriklausomų imčių t-testą tarp grupių, kurios matė įrašus su DI sąsaja (grupės 1 ir 2), ir grupių, kurios matė įrašus be DI sąsajos (grupės 3 ir 4), buvo gauti šie rezultatai: T statistika = 2.36 ir $p = 0.0187$, rodantys statistiškai reikšmingą skirtumą. Nors šis klausimas buvo įtrauktas į tyrimo klausimyną kaip priemonė manipuliacijos efektyvumui įvertinti, jis nebuvo įtrauktas į galutinį instrumentą, nes statistiškai reikšmingas skirtumas tarp grupių ($p = 0.0187$) parodė, kad dalyviai aiškiai ir nedviprasmiškai suvokia turinio ryšį su dirbtiniu intelektu. Vis dėlto, naudojant didelius kalbos modelius tyrimo instrumento peržiūrai, būtina atsižvelgti į tam tikrus apribojimus. Pirma, kaip pastebi Sarstedt et al. (2024), dideli kalbos modeliai negali visiškai atkartoti realių žmonių atsakymų, ypač kai kalbama apie subjektyvius vertinimus. Antra, modelio išvados turi būti traktuojamos kaip pagalbinės rekomendacijos, o ne galutiniai vertinimai. Dėl šių priežasčių šiame tyrime „ChatGPT“ sugeneruoti duomenys buvo naudojami išskirtinai tik tyrimo scenarijų ir matavimo priemonių struktūros testavimui – siekiant įvertinti manipuliacinių sąlygų loginį nuoseklumą. Šie duomenys nebuvo įtraukti į galutinę analizę, o jų paskirtis buvo metodologinė – užtikrinti, kad eksperimento dizainas būtų aiškus, pagrįstas ir gebėtų diferencijuoti tarp eksperimentinių grupių dar prieš renkant tikruosius dalyvių atsakymus.

Šiame tyrime „ChatGPT“ taip pat buvo pasitelktas kaip pagalbinė priemonė scenarijais grįstam eksperimentui parengti, siekiant išanalizuoti DI sukurtų socialinių tinklų pranešimų žymėjimo poveikį vartotojų suvokiamam šaltinio ir turinio patikimumui bei įsitraukimui. Kadangi tyrimas yra grįstas eksperimento dizainu ir scenarijais, būtina buvo sukurti du lygiaverčius, tačiau skirtingai pažymėtus turinio vienetų, kurie atspindėtų realistišką ir aktualų mokslinės inovacijos pristatymą.

Tam pasirinktas „ChatGPT“ įrankis dėl jo gebėjimo generuoti rišlius, stilistiškai pritaikytus ir kontekstualiai atitinkančius tekstus pagal tikslias instrukcijas.

Pasirinkimas naudoti viešai prieinamus Kauno technologijos universiteto (KTU) pranešimus, publikuotus žiniasklaidoje (15min.lt ir [Delfi.lt](https://delfi.lt)), buvo pagrįstas siekiu remtis aktualia, patikrinta, profesionaliai suformuota ir Lietuvos kontekstui būdinga informacija. Šie pranešimai aptaria realiai vykdomus mokslo projektus bei inovacijas, todėl leido užtikrinti eksperimento turinio pagrįstumą ir tikroviškumą. Be to, universitetų komunikacija spaudai paprastai būna adaptuota plačiajai auditorijai, tačiau išlaiko pakankamą informatyvumo lygį, todėl tai tapo tinkamu pagrindu turiniui, kuris vėliau buvo perdirbtas tiek akademiniam, tiek viešam vartojimui. Instrukcijos, pateiktos „ChatGPT“ (žr. 5 priedas), buvo suformuluotos taip, kad užtikrintų nuoseklią komunikacijos strategiją – sukurti glaustus, suprantamus, emociniu tonu neišsiskiriančius socialinių tinklų pranešimus. Pastarieji buvo parengti informaciniam „Facebook“ įrašui, kuris turėjo būti vizualiai, stilistiškai ir kalbiniu tonu pritaikytas plačiajai auditorijai. Tokia dviguba turinio forma atliepia tyrimo poreikį testuoti vartotojų reakcijas į DI generuatą komunikaciją skirtinguose kontekstuose.

Svarbu pažymėti, kad tyrime taikytas scenarijais grįstas eksperimento dizainas, kuris socialinių mokslų tyrimuose yra dažnai pasitelkiamas analizuojant priežastinius ryšius tarp kintamųjų. Šiuo atveju buvo siekiama išsiaiškinti, kaip DI sukurtas turinys, pažymėtas atitinkama etikete (nurodančia, kad tekstą sukūrė DI), veikia auditorijos požiūrį, lyginant su analogišku, bet nepažymėtu turiniu. „ChatGPT“ pasitelkimas leido sukurti vienodai struktūruotus, stilistiškai panašius, tačiau eksperimentui pritaikytus skirtingus scenarijus, taip užtikrinant tyrimo vidinį validumą.

Tyrimo vykdymas. Duomenų rinkimas buvo vykdomas taikant netikimybinę patogiąją atranką, nuo 2025 m. balandžio 6 d. iki 2025 m. balandžio 21 d. Apklausą buvo platinama mišriu būdu, siekiant pasiekti kuo įvairesnę respondentų auditoriją ir užtikrinti pakankamą tyrimo imties dydį. Pirmiausia apklausos nuoroda buvo viešai paskelbta keliose specializuotose socialinio tinklo „Facebook“ grupėse, skirtose tyrimams ir apklausoms („APKLAUSOS“, „Apklausos!“), taip pat asmeninėse „Facebook“, „Instagram“, „LinkedIn“ socialinių tinklų platformų paskyrose, kviečiant dalyvauti tyrime platesnę socialinių tinklų naudotojų bendruomenę. Papildomai apklausą buvo išsiųsta Kauno technologijos universiteto studentams. Ji buvo įtraukta į kassavaitinį universiteto naujienlaiškį, taip sudarant sąlygas į tyrimą įsitraukti aukštosios mokyklos bendruomenei.

Duomenų analizės procedūros. Surinkti tyrimo duomenys buvo analizuojami naudojantis statistinės analizės programine įranga „IBM SPSS Statistics 29.0.2.0“. Kadangi dalis semantinio diferencialo skalės teiginių buvo formuluoti skirtingomis kryptimis, prieš atliekant analizę šių teiginių atsakymai buvo atvirkščiai perkoduoti. Pradinėje analizės stadijoje buvo atlikta dažnių analizė, siekiant apibūdinti pagrindines tiriamųjų charakteristikas. Analizuoti šie rodikliai: amžius, lytis, išsilavinimas, mėnesinės pajamos po mokesčių, socialinių tinklų naudojimo dažnis ir dažniausiai naudojamos platformos. Tolesniame etape buvo įvertinta tyrimo instrumento metodologinė kokybė. Skalių struktūros tinkamumas buvo tikrinamas naudojant Kaiser-Meyer-Olkin (KMO) imties adekvatumo matą bei Bartlerto sferiškumo kriterijų. Siekiant patikrinti, ar pasirinkti teiginiai tinkamai atspindi nagrinėjamus konstruktus, buvo atlikta faktorinė analizė, pasitelkiant pagrindinių komponentų analizės (PCA) ir „Direct Oblimin“ faktorių sukimo metodus. „Direct Oblimin“ metodas buvo pasirinktas dėl to, jog jis leidžia atsižvelgti į galimą faktorių tarpusavio koreliaciją – galima gauti tikslesnį ir empiriškai labiau pagrįstą faktorių išsidėstymą, kuris atspindi realų tarpusavio ryšį tarp latentinių kintamųjų. Pastarojo meto metodologinėje literatūroje rekomenduojama taikyti „Direct

Oblimin“ pasukimo metodą (Howard, 2023). Skalių patikimumui įvertinti buvo apskaičiuotos Kronbacho alfa (angl. *Cronbach Alfa*) reikšmės kiekvienam tyrimo konstruktui. Vėliau atlikta aprašomoji analizė, kurioje pateikiamos pagrindinės priklausomųjų kintamųjų (suvokiamo šaltinio ir turinio atitikimo, šaltinio patikimumo, turinio patikimumo bei įsitraukimo ketinimų) statistinės charakteristikos: minimali ir maksimali reikšmės, vidurkis bei standartinis nuokrypis. Tyrime taip pat buvo taikytos neparametrinės analizės procedūros, tokios kaip Kruskal–Wallis testai, skirtumams tarp grupių (eksperimentinių scenarijų) vertinti, kai kintamųjų pasiskirstymas neatitiko normalumo prielaidų. Normalumo testavimui buvo naudojamas Kolmogorovo–Smirnov (K–S) testas. Priklausomai nuo normalumo testų rezultatų, koreliacijos analizei tarp tiriamųjų kintamųjų buvo taikytas Spearmano koreliacijos koeficientas. Papildomai tyrimo duomenims analizuoti buvo atlikta ir kokybinė turinio analizė, pasitelkiant sentimentų analizės metodą. Šis metodas buvo taikytas atviriems tyrimo dalyvių atsakymams analizuoti, siekiant identifikuoti dominuojančias emocines nuostatas DI sukurto turinio atžvilgiu. Galiausiai, tyrimo hipotezėms tikrinti buvo taikytos tiek dispersinės analizės (ANOVA), tiek tiesinės regresinės analizės procedūros. Dispersinė analizė leido įvertinti eksperimentinių manipuliacijų – skirtingų scenarijų – poveikį priklausomiesiems kintamiesiems, o regresinės analizės buvo pasitelktos siekiant detaliau ištirti prognozinį ryšius tarp tyrime naudotų konstrukto.

4. Dirbtinio intelekto sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose, empirinio tyrimo rezultatai

Šiame magistro baigiamojo darbo skyriuje pristatomi ir interpretuojami empirinio tyrimo metu surinkti duomenys, atliekamas hipotezių patikrinimas, o taip pat aptariami gauti rezultatai platesniame kontekste. Skyriuje išskiriamos galimos tyrimo rezultatų praktinio panaudojimo kryptys, tyrimo ribotumai bei siūlomos potencialios tolimesnių mokslinių tyrimų kryptys.

4.1. Tyrimo imties charakterizavimas pagal sociodemografinės charakteristikos

Galutinėje imtyje likusių 372 respondentų, dalyvavusių eksperimento dizaino tyrime, duomenys pirmiausia analizuojami pagal bendrąsias charakteristikas. 5 lentelėje pateikiami pagrindiniai demografiniai rodikliai – lytis, amžius, išsilavinimas bei mėnesinės pajamos po mokesčių. Šie duomenys leidžia apibūdinti tyrimo imties struktūrą ir įvertinti jos atitiktį tyrimo tikslams. Išsamesnė demografinių rodiklių analizė pateikiama 6 priede.

5 lentelė. Tyrimo respondentų demografinės charakteristikos (N = 372)

Demografiniai duomenys		Respondentų pasiskirstymas (N)	Procentinė dalis (proc.)
Lytis	Vyras	150	40,3
	Moteris	220	59,1
	Kita	2	0,5
Amžius	15–23 m.	105	28,2
	24–25 m.	81	21,8
	26–31 m.	96	25,8
	32–67 m.	90	24,2
Išsilavinimas	Pradinis išsilavinimas	62	16,7
	Pagrindinis išsilavinimas	24	6,5
	Vidurinis išsilavinimas	212	57,0
	Aukštesnysis arba specialusis vidurinis išsilavinimas	6	1,6
	Aukštasis išsilavinimas	68	18,3
Pajamos	Mažiau nei 500 Eur	26	7,0
	500–999 Eur	50	13,4
	1000–1499 Eur	114	30,6
	1500–1999 Eur	68	18,3
	2000–2499 Eur	26	7,0
	2500 Eur ar daugiau	30	8,1
	Nenoriu atsakyti	58	15,6

Remiantis tyrimo rezultatais, pagal lytinę tapatybę, didesnė dalis apklausos dalyvių buvo moterys – jos sudarė 59,1 proc. visos imties, o vyrų dalis siekė 40,3 proc. Kitos lyties asmenų buvo 0,5 proc. Amžius buvo matuojamas santyktine skale, todėl buvo apskaičiuotas aritmetinis vidurkis, kuris siekė 29,44 metų (SD – 10,08). Nors vidutinė reikšmė rodo, kad imtyje vyrauja jaunesnio amžiaus respondentai, mediana (25,5 m.) ir moda (25 m.) buvo mažesnės nei vidurkis, o tai leidžia teigti, jog skirstinys yra šiek tiek asimetriškas – t. y. pasitaikė ir vyresnių respondentų, kurie galėjo kilstelėti

bendrą vidurkį. Atsižvelgiant į tai, kad amžiaus kintamasis neatitiko normaliojo skirstinio, jis buvo transformuotas į keturias kategorines grupes pagal kvartilines reikšmes. Skirstymui naudotos 25-ojo, 50-ojo ir 75-ojo procentilių ribos (atitinkamai 23, 25,5 ir 31 metai), todėl buvo suformuotos tokios amžiaus grupės: 15–23 m., 24–25 m., 26–31 m. ir 32–67 m..

Didžiausią dalį sudarė jauniausi tyrimo dalyviai – 28,2 proc. priklausė 15–23 metų amžiaus grupei. Antroje vietoje pagal dydį buvo 26–31 metų grupė (25,8 proc.), kurią nežymiai lenkė vyriausių respondentų grupė – 32–67 metų (24,2 proc.). Mažiausią dalį sudarė 24–25 metų grupė, kuri apėmė 21,8 proc. visos imties. Toks pasiskirstymas rodo pakankamai tolygų imties išsidėstymą tarp visų amžiaus grupių, leidžiantį atlikti tarpusavio palyginimus.

Pagal išsilavinimą, didžiausią dalį tyrimo imtyje sudarė vidurinį išsilavinimą nurodę respondentai – jie sudarė 57 proc. visų dalyvių. Atsižvelgiant į tai, kad beveik 60 proc. respondentų priklausė 15–25 metų amžiaus grupėms, tikėtina, kad dalis jų dar nebuvo įgiję aukštojo išsilavinimo tyrimo metu. Kitos išsilavinimo kategorijos buvo mažiau gausios: pradinį išsilavinimą nurodė 16,7 proc., aukštąjį – 18,3 proc., o pagrindinį bei aukštesnįjį ar specialųjį vidurinį – atitinkamai 6,5 proc. ir 1,6 proc. Nors bendra išsilavinimo struktūra atspindi tam tikrą įvairovę, didelė vidurinio išsilavinimo dalis rodo, kad tyrimo metu imtyje dominavo jaunesnio amžiaus, dar studijuojantys ar neseniai mokyklą baigę asmenys. Tai atitinka įprastą demografinį pasiskirstymą tarp aktyviausių vartotojų populiariausiose socialinių tinklų platformose pasaulyje, kuriose dominuoja 18–24 metų amžiaus grupė (Dixon, 2024c, 2024d; Ceci, 2025).

Vertinant respondentų finansinę padėtį pagal mėnesines pajamas po mokesčių, daugiausia respondentų nurodė gaunantys 1000–1499 Eur (30,6 proc.) ir 1500–1999 Eur (18,3 proc.). Kiek mažesnė dalis – po 7,0 proc. – įvardijo gaunantys mažiau nei 500 Eur arba 2000–2499 Eur, o 8,1 proc. nurodė 2500 Eur ar daugiau. Pažymėtina, kad 15,6 proc. respondentų pasirinko neatskleisti savo pajamų. Toks atsakymų pasiskirstymas rodo, kad dauguma tyrimo dalyvių patenka į vidutinių pajamų kategoriją, o tai leidžia analizuoti vartotojų reakcijas ne tik pagal amžių ar išsilavinimą, bet ir pagal ekonominę padėtį.

Empirinio tyrimo metu respondentų buvo prašoma nurodyti, kiek vidutiniškai laiko per dieną jie praleidžia socialiniuose tinkluose. Šių atsakymų pasiskirstymas pateiktas 6 lentelėje.

6 lentelė. Respondentų pasiskirstymas pagal socialinių tinklų naudojimo trukmę per dieną (N = 372)

Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose?	Pasirinko (N)	Procentas (%)
1 valandą ir mažiau	54	14,5
2 valandas	118	31,7
3 valandas	114	30,6
4 valandas	52	14,0
5 valandas ir daugiau	34	9,1

Gauti rezultatai atskleidė, kad didžioji dalis respondentų kasdien socialiniuose tinkluose praleidžia nuo dviejų iki trijų valandų. Būtent 31,7 proc. tyrimo dalyvių nurodė, kad socialiniuose tinkluose praleidžia vidutiniškai dvi valandas per dieną, o dar 30,6 proc. – tris valandas. Šie du atsakymų variantai sudaro beveik du trečdalius visų atsakymų (62,3 proc.). Tuo tarpu 14,5 proc. respondentų teigė, kad socialiniams tinklams skiria tik vieną valandą ar mažiau per dieną, o 14,0 proc. – keturias valandas. Palyginti nedidelė tyrimo dalis – 9,1 proc. – nurodė, kad kasdien socialiniams tinklams

skiria penkias valandas ar daugiau. Šie rezultatai leidžia daryti prielaidą, kad dauguma respondentų socialinius tinklus naudoja reguliariai, tačiau jų naudojimo intensyvumas išlieka vidutinio lygmens.

Tyrimo metu respondentų taip pat buvo prašoma pažymėti, kurias socialinių tinklų platformas jie naudoja dažniausiai. Rezultatai, pateikti 7 lentelėje, rodo, kad labiausiai paplitę socialiniai tinklai tarp apklaustųjų buvo „Instagram“ ir „Facebook“ – atitinkamai šias platformas pasirinko 79,6 proc. ir 77,4 proc. visų tyrimo dalyvių. Trečioje vietoje pagal populiarumą buvo „YouTube“, kurį naudojo 57 proc. respondentų.

7 lentelė. Respondentų dažniausiai naudojami socialiniai tinklai (N = 372)

Socialinis tinklas	Pasirinko (N)	Procentas (%)
„Instagram“	296	79,6 %
„Facebook“	288	77,4 %
„YouTube“	212	57,0 %
„TikTok“	128	34,4 %
„LinkedIn“	82	22,0 %
„Telegram“	36	9,7 %
„X (Twitter)“	22	5,9 %
„Reddit“	14	3,8 %
„Pinterest“	4	1,1 %
„Spotify“	4	1,1 %
„Discord“	2	0,5 %
„Mastodon“	2	0,5 %

Reikšmingą, tačiau mažesnę respondentų dalį sudarė „TikTok“ vartotojai – šią platformą nurodė 34,4 proc. dalyvių. Nors ši platforma dažniau siejama su jaunesnio amžiaus auditorijomis, tyrimo rezultatai rodo jos reikšmingą įsitvirtinimą tarp platesnės vartotojų grupės. Tuo tarpu „LinkedIn“ pasirinko 22,0 proc. respondentų, o tai gali atspindėti šios platformos profesinę specifiką ir siauresnį naudojimo kontekstą. Mažesnio populiarumo susilaukė ir tokios platformos kaip „Telegram“ (9,7 proc.), „X (Twitter)“ (5,9 proc.), „Reddit“ (3,8 proc.), „Pinterest“ ir „Spotify“ (po 1,1 proc.), bei „Discord“ ir „Mastodon“ – šias pažymėjo vos po 0,5 proc. tyrimo dalyvių. Šie rezultatai leidžia daryti prielaidą, kad pastarosios platformos yra labiau nišinės ir aktualios tik specifinėms vartotojų grupėms.

Apibendrinant galima teigti, kad tyrimo imtis pasižymėjo tam tikru amžiaus, lyties ir socioekonominio statuso pasiskirstymo balansu, tačiau išsiskyrė jaunesnio amžiaus ir moterų dominavimu, o tai vertinant socialinių tinklų vartotojų profilius yra aktualu ir atitinka šiuolaikinių komunikacijos praktikų realijas. Taip pat, tyrimo imtis ne tik atspindi dominuojančias socialinių tinklų naudojimo tendencijas, bet ir pagrindžia tyrime taikytos manipuliacijos kontekstinę aktualumą. Atsižvelgiant į tai, kad net 77,4 proc. respondentų nurodė dažniausiai naudojantys „Facebook“, šios platformos pasirinkimas manipuliaciniams eksperimentiniams scenarijams imituoti buvo metodiškai pagrįstas. Tyrimo metu visiems dalyviams buvo pateikiami socialinių tinklų įrašai, stilistiškai ir vizualiai modeliuoti kaip „Facebook“ publikacijos. Tokia tyrimo konstrukcija užtikrina, kad eksperimentiniai vaizdai būtų atpažįstami, realistiški ir kontekstualiai susieti su dalyvių kasdieniu medijų vartojimu. Tai ypač svarbu, vertinant DI žymėjimo poveikį turinio autorystės suvokimui,

kadangi tyrime naudotas turinio pateikimas imitavo natūralią skaitmeninę aplinką, kurioje vartotojai įprastai priima sprendimus dėl turinio autorystės ir atsako į tai.

4.2. Tyrimo instrumento metodologinės kokybės analizė

Siekiant įvertinti tiriamųjų konstrukto struktūrą – ar pasirinkti teiginiai iš tiesų atspindi vieną arba kelis teorinius matmenis, kuriuos norima tirti – buvo pasitelkti du pagrindiniai rodikliai – Kaiser–Meyer–Olkin (KMO) imties adekvatumo matas ir Bartleto sferiškumo kriterijus (Pallant, 2020) (žr. 7 priedas). KMO reikšmės visoms skalėms yra $\geq 0,6$, o tai, remiantis Kaiser'io (1974) klasifikacija, laikoma tinkamu rezultatu faktorių analizei (žr. 8 lentelė). Viena iš skalių – šaltinio patikimumo – pasižymi itin aukšta KMO reikšme (0,911), kas rodo labai aukštą kintamųjų tarpusavio ryšių tinkamumą. Kitos skalės, nors ir turinčios žemesnes KMO reikšmes (nuo 0,611 iki 0,742), taip pat atitinka minimalų kriterijų (0,6) (Pallant, 2020), leidžiantį atlikti faktorių analizę. Vertėtų pažymėti, kad šaltinio ir turinio atitikimo bei įsitraukimo skalėse teiginių skaičius buvo nedidelis (po 3), todėl ir KMO reikšmės natūraliai žemesnės, nes šis rodiklis yra jautrus kintamųjų kiekiui.

8 lentelė. Skalių tinkamumo vertinimas tiriamajai faktorių analizei (N = 372)

Skalės pavadinimas	Teiginių skaičius	KMO imties adekvatumo matas	Bartleto sferiškumo kriterijaus p reikšmė
Šaltinio ir turinio atitikimo skalė	3	0,611	<0,001
Turinio patikimumo skalė	3	0,742	<0,001
Šaltinio patikimumo skalė	18	0,911	<0,001
Vartotojų įsitraukimo ketinimų skalė	3	0,681	<0,001

Visose skalėse Bartleto sferiškumo testas buvo statistiškai reikšmingas ($p < 0,001$), o tai reiškia, kad tarp skalės teiginių egzistuoja pakankamai stiprūs ryšiai, leidžiantys pagrįstai taikyti faktorių analizę.

Siekiant empiriškai įvertinti, ar tyrime naudojami teiginiai logiškai grupuojasi į atskirus konstrukcinius vienetus ir ar šie vienetai atitinka conceptualų tyrimo modelį, buvo atlikta tiriamoji faktorių analizė. Analizė buvo vykdyta visiems konstrukto atskirai, naudojant pagrindinių komponentų išskyrimo metodą (angl. *Principal Component Analysis*) ir dažniausiai tyrimuose naudojamą „Direct Oblimin“ faktorių sukimą (Pallant, 2020), kuris leidžia faktoriams koreliuoti. Tokia metodika pasirinkta remiantis prielaida, kad vartotojų požiūrio komponentai gali būti tarpusavyje susiję. Gauti rezultatai atskleidė penkis aiškius faktorius, kurių kiekvienas atitiko atskirą vartotojų suvokimo ar atsako dimensiją. Teiginiai buvo priskirti faktoriams pagal jų apkrovų dydžius (angl. *factor loadings*) bei remiantis aiškiu teoriniu pagrindu – tyrime naudotos skalės buvo sukurtos atsižvelgiant į jau egzistuojančius mokslinių tyrimų instrumentus, kurie anksčiau buvo empiriškai pagrįsti kitų autorių darbais. Papildomai atsižvelgta į tai, ar konkretūs teiginiai pasižymėjo vienareikšmiu priskyrimu vienam faktoriui. Teiginiai, kurių svoriai reikšmingai persidengė tarp kelių komponentų ar kurių semantinis turinys neatitiko aiškos konstrukto struktūros, buvo pašalinti iš galutinio modelio. Tokiu būdu išgryninta konstrukto struktūra ir užtikrintas rezultatuose pagrįstas komponentų aiškumas, leidžiantis toliau modeliuoti šaltinio patikimumo struktūrą kaip sudėtinį, daugiamačių reiškinį.

Atlikus tiriamąją faktorių analizę su šaltinio patikimumo skalės teiginiais (angl. *source credibility*) (McCroskey & Teven, 1999), paaiškėjo, kad konstrukto struktūra neatitiko pradinės, vienalytės skalės

sampratos. Nors teoriniu požiūriu šaltinio patikimumas dažniausiai yra operacionalizuojamas kaip vienas konstruktas, jis taip pat gali būti skaidomas į tris pagrindines dimensijas: kompetenciją, palankumą ir pasitikėjimo vertumą, kaip tai numatė McCroskey'us ir Teven'as (1999). Tačiau šiame tyrime atlikta faktorinė analizė atskleidė, kad respondentai skirtingus šaltinio aspektus suvokė kaip iš esmės atskirus požymius. Sprendžiant pagal komponentų svorius ir jų pasiskirstymą, šaltinio patikimumas išsiskaidė į keturias komponentes, kurių kiekviena atspindinti tam tikrą šaltinio vertinimo pobūdį (žr. 9 lentelė). Šių dimensijų pavadinimai buvo suformuoti ieškant prasminių sąsajų tarp teiginių, sudarančių kiekvieną komponentą, siekiant konceptualiai įprasminti jų turinį.

9 lentelė. Šaltinio patikimumo konstrukto teiginiai pagal išskirtus faktorius (N = 372)

Skale matuojamas konstruktas	Skalės teiginiai	Išskirtas faktorius (komponentė) 1	Išskirtas faktorius (komponentė) 2	Išskirtas faktorius (komponentė) 3	Išskirtas faktorius (komponentė) 4
Šaltinio patikimumas per kompetenciją	[Ne ekspertas – Ekspertas]	0.854			
	[Neapmokytas – Apmokytas]	0.791			
	[Nekompetentingas – Kompetentingas]	0.765			
	[Protingas – Neprotingas] (R)	0.624			
	[Sumanus – Kvailas] (R)	0.576			
	[Informuotas – Neinformuotas] (R)	0.474			
Šaltinio patikimumas per palankumą	[Domisi manimi – Nesidomi manimi] (R)		0.908		
	[Rūpinasi manimi – Nesirūpina manimi] (R)		0.859		
	[Rūpinasi mano interesais – Nesirūpina mano interesais] (R)		0.824		
Šaltinio patikimumas per empatiją	[Savanaudis – Nesavanaudis]			0.775	
	[Nejautrus – Jautrus]			0.730	
	[Nesupratingas – Supratingas]			0.676	
Šaltinio patikimumas per pasitikėjimą	[Garbingas – Negarbingas] (R)				–0.851
	[Sąžiningas – Nesąžiningas] (R)				–0.840
	[Moralus – Amoralus] (R)				–0.797

Pirmasis faktorius atliepė **šaltinio patikimumo per kompetenciją** dimensiją. Jam priskirti teiginiai, susiję su autoriaus gebėjimais ir išmanymu, tokie kaip „kompetentingas“, „ekspertas“, „informuotas“. Svoriai šioje grupėje svyravo nuo 0,474 iki 0,854, o tai rodo aiškų konstrukto atitikimą ir stiprų vidinį nuoseklumą.

Antrasis faktorius išskyrė **šaltinio patikimumą per palankumą**, įtraukiant teiginius, susijusius su tuo, ar autorius rodo rūpestį ir asmeninį dėmesį. Stipriausius svorius turintys elementai buvo „rūpinasi manimi“, „domisi manimi“, kurių svoriai viršijo 0,8.

Trečiasis faktorius atspindėjo **šaltinio patikimumo per empatiją** dimensiją, kuriai priskirti tokie teiginiai kaip „nesavanaudis“ (0,775), „jautrus“ (0,730) ir „supratingas“ (0,676). Tai rodo, kad dalis respondentų vertina šaltinio patikimumą per žmogiškumo ir emocinio jautrumo prizmę.

Ketvirtasis faktorius atskyrė **šaltinio patikimumo per pasitikėjimą** požymius – jam priskirti tokie teiginiai kaip „garbingas“, „sąžiningas“ ir „moralus“. Nors svoriai buvo neigiami, jų reikšmės buvo labai aukštos (nuo –0,797 iki –0,851) – rodo aiškiai apibrėžtą atskirą dimensiją. Atlikus tiriamąją faktoriinę analizę, paaiškėjo, kad šios šaltinio patikimumo konstrukto dimensijos paaiškina 75,132 proc. bendros duomenų dispersijos (žr. 7 priedas).

Toliau buvo atlikta vartotojų įsitraukimo ketinimų skalės tiriamoji faktoriinė analizė. Šiam faktoriui buvo priskirti trys teiginiai: apie pasidalijimą įrašu, komentavimą ir paspaudimą „patinka“. Visi šie teiginiai buvo glaudžiai susiję tarpusavyje ir stipriai siejosi su ta pačia komponentę – jų faktoriniai svoriai siekė nuo 0,834 iki 0,909 (žr. 10 lentelė), o teiginiai paaiškino 75,132 proc. bendros duomenų dispersijos (žr. 8 priedas).

10 lentelė. Vartotojų įsitraukimo ketinimų konstrukto teiginiai pagal išskirtus faktorius (N = 372)

Skale matuojamas konstruktas	Skalės teiginiai	Išskirtas faktorius (komponentė) 1
Vartotojų įsitraukimo ketinimai	Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, pasidalinčiau juo su kitais	0,909
	Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, palikčiau komentarą po šiuo įrašu	0,843
	Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, paspausčiau „patinka“ po šiuo įrašu	0,834

Toliau buvo tikrinama šaltinio ir turinio atitikimo suvokimo (angl. *source-message congruence*) (Choi, 2024; Lee et al., 2025) skalė. Į vieną išskirtą faktorių pateko teiginiai, kuriuose vertinamas semantinis ryšys tarp turinio ir jo autoriaus, pavyzdžiui: „įrašas ir jo autorius atitinka vienas kitą“, „autorius veiksmingai perteikė numatytą žinutę“. Svoriai šiame faktoriuje varijavo nuo 0,716 iki 0,882 (žr. 11 lentelė), patvirtindami konstrukto empirinį pagrįstumą ir conceptualinį nuoseklumą. Šaltinio ir turinio atitikimo konstrukto teiginiai paaiškino 64,751 proc. visos dispersijos (žr. 9 priedas).

11 lentelė. Suvokiamo šaltinio ir turinio atitikimo konstrukto teiginiai pagal išskirtus faktorius (N = 372)

Skale matuojamas konstruktas	Skalės teiginiai	Išskirtas faktorius (komponentė) 1
Suvokiamas šaltinio ir turinio atitikimas	Manau, kad šis socialinių tinklų įrašas ir jo autorius atitinka vienas kitą	0,882

	Manau, kad šis socialinių tinklų įrašas yra susijęs su jo autoriumi	0,808
	Manau, kad socialinių tinklų įrašo autorius veiksmingai perteikė numatytą žinutę	0,716

Toliau buvo analizuojama turinio patikimumo skalė (angl. *message credibility*) (Brüns & Meißner, 2024), apimanti tokias savybes kaip įrašo įtikinamumą, sąžiningumą ir tai, ar įrašas kelia pasitikėjimą. Visi trys skalės teiginiai formavo vieną aiškų komponentą, kuris paaiškino 83,37 proc. visos duomenų dispersijos (žr. 10 priedas). Tai yra labai aukštas paaiškinamos dispersijos procentas, leidžiantis pagrįstai teigti, kad šie teiginiai matuoja vientisą, nuoseklų konstrukto turinį. Taip pat, faktoriui priskirti teiginiai turėjo aukštus svorius – nuo 0,892 iki 0,929 (žr. 12 lentelė).

12 lentelė. Suvokiamo turinio patikimumo konstrukto teiginiai pagal išskirtus faktorius (N = 372)

Skale matuojamas konstruktas	Skalės teiginiai	Išskirtas faktorius (komponentė) 1
Suvokiamas turinio patikimumas	Šis socialinių tinklų įrašas atrodo įtikinamas	0,929
	Šis socialinių tinklų įrašas kelia pasitikėjimą	0,918
	Šis socialinių tinklų įrašas atrodo sąžiningas	0,892

Siekiant įvertinti tiriamųjų konstrukto matavimo skalių vidinį nuoseklumą ir tinkamumą tolesnei statistinei analizei, buvo apskaičiuoti Kronbacho alfa koeficientai kiekvienai konstrukto skalei atskirai (žr. 13 lentelė ir 11, 12, 13, 14, 15, 16 ir 17 prieduose). Šis rodiklis leidžia įvertinti, ar skalė pasižymi vidiniu suderinamumu (Pallant, 2010).

13 lentelė. Tiriamų konstrukto skalių patikimumo vertinimas (N = 372)

Skalės pavadinimas	Dimensijos	Teiginių skaičius	Kronbacho alfa koeficientas
Turinio patikimumas (žr. 12 priedas)		3	0,900
Šaltinio patikimumas	Pasitikėjimas šaltiniu (žr. 11 priedas) (3)	18	0,903
	Šaltinio kompetencija (žr. 13 priedas) (6)		0,889
	Šaltinio palankumas (žr. 14 priedas) (3)		0,855
	Šaltinio empatija (žr. 15 priedas) (3)		0,828
Įsitraukimo ketinimai (žr. 16 priedas)		3	0,820
Šaltinio ir turinio atitikimas (žr. 17 priedas)		3	0,722

Rezultatai rodo, jog visų septynių skalių Kronbacho alfa reikšmės viršijo 0,7 ribą, kuri laikoma minimalia priimtina patikimumo norma (Nunnally, 1978). Aukščiausią patikimumą demonstravo šaltinio pasitikėjimo ($\alpha=0,903$) ir turinio patikimumo ($\alpha=0,900$) skalės. Kiek žemesni, tačiau vis tiek aukšti patikimumo rodikliai buvo fiksuoti šaltinio kompetencijos ($\alpha=0,889$), šaltinio palankumo ($\alpha=0,855$), šaltinio empatijos ($\alpha=0,828$) ir įsitraukimo ketinimų ($\alpha=0,820$) skalėse. Šie rezultatai leidžia teigti, kad visos tyrime naudotos konstrukto skalės yra patikimos ir tinkamos naudoti tolimesniuose analizės etapuose.

Atlikus tiriamąją faktoriinę analizę su šaltinio patikimumo skalės teiginiais, pastebėta, jog šio konstrukto struktūra išsidėsto į keturias atskiras dimensijas: kompetenciją, palankumą, empatiją ir pasitikėjimą. Dėl šios priežasties tikslinamos H3, H6 ir H8 hipotezės, kuriomis siekiama nustatyti tiesioginį šaltinio patikimumo poveikį vartotojo įsitraukimo ketinimams socialiniuose tinkluose bei įvertinti netiesioginį šaltinio ir turinio atitikimo poveikį vartotojo įsitraukimo ketinimams per įsiterpiančią šaltinio patikimumo konstruklą (ir atskiras jo dimensijas). Žemiau pateikiamas pakoreguotų hipotezių sąrašas, kurios bus tikrinamos empirinio tyrimo metu:

H3a: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų suvokiamą šaltinio kompetenciją.

H3b: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų suvokiamą šaltinio empatiją.

H3c: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų suvokiamą šaltinio palankumą.

H3d: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų pasitikėjimą šaltiniu.

H6a: Vartotojų suvokiama šaltinio kompetencija teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;

H6b: Vartotojų suvokiama šaltinio empatija teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;

H6c: Vartotojų suvokiamas šaltinio palankumas teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;

H6d: Vartotojų pasitikėjimas šaltiniu teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;

H8a: Šaltinio kompetencija teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose;

H8b: Šaltinio empatija teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose;

H8c: Šaltinio palankumas teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose;

H8d: Pasitikėjimas šaltiniu teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose.

Kitos hipotezės, nesusijusios su šaltinio patikimumo dimensijomis, lieka nekeičiamos.

4.3. Vartotojų atsako į DI sukurtą turinį skalių aprašomoji analizė

Siekiant įvertinti galimus netipinius atvejus ir išsiskiriančius tyrimo kintamųjų skalių elementus, buvo atlikta aprašomoji analizė (žr. 18 priedas). Atliekant šią analizę, buvo įvertintos kiekvieno konstrukto – šaltinio patikimumo per pasitikėjimą, kompetenciją, palankumą, empatiją, taip pat turinio

patikimumo, šaltinio ir turinio atitikimo bei įsitraukimo ketinimų – skalės. 14 lentelėje pateikiamos kiekvienos šių skalių minimali ir maksimali reikšmės, vidurkiai bei standartiniai nuokrypiai, leidžiantys preliminariai apibūdinti vartotojų nuostatas ir jų pasiskirstymą.

14 lentelė. Tyrimo kintamųjų vertinimo charakteristikos (N = 372)

Skalės pavadinimas	Dimensijos	Minimali reikšmė	Maksimali reikšmė	Vidurkis	Standartinis nuokrypis
Šaltinio ir turinio atitikimas		1,33	7,00	4,3441	1,12788
Turinio patikimumas		1,00	7,00	4,1846	1,39909
Įsitraukimo ketinimai		1,00	7,00	2,6326	1,35142
Šaltinio patikimumas	Šaltinio kompetencija (6)	1,00	7,00	4,7590	1,24337
	Šaltinio palankumas (3)	1,00	7,00	3,9749	1,41674
	Šaltinio empatija (3)	1,00	7,00	4,5036	1,26758
	Pasitikėjimas šaltiniu (3)	1,00	7,00	4,6649	1,36464

Rezultatai rodo, kad aukščiausiai buvo įvertintas šaltinio patikimumas per pasitikėjimą ($M = 4,66$; $SD = 1,36$) bei šaltinio kompetencija ($M = 4,76$; $SD = 1,24$) – tyrimo respondentai DI kaip įrašo autorių vertina kaip labiau garbingą, sąžiningą bei kompetentingą. Tuo tarpu mažiausią vidurkį pasiekė vartotojų įsitraukimo ketinimų vertinimas ($M = 2,63$; $SD = 1,35$), rodantis atsargesnį respondentų elgseninį atsaką – tikėtina, jog respondentai nėra linkę dalintis DI sukurtu turiniu, jį komentuoti ar spausti mygtuką „patinka“ socialiniuose tinkluose. Aukščiausią minimalią reikšmę tarp visų konstrukto vidurkių turėjo „Šaltinio ir turinio atitikimo“ skalė ($Min = 1,33$; $M = 4,34$; $SD = 1,13$). Tai reiškia, kad nė vienas iš respondentų nepateikė žemiausio galimo įvertinimo – 1,00. Šį skirtumą paaiškina tai, jog skalę sudarė keli teiginiai, kurių reikšmės buvo sujungtos apskaičiuojant vidurkį, todėl apatinė skalės reikšmė susiformavo iš kelių žemų, bet ne pačių žemiausių, vertinimų. Kadangi nė vienas respondentų visų trijų teiginių neįvertino mažiausiu balu, bendra skalės minimali reikšmė natūraliai buvo aukštesnė nei 1,00. Toks rezultatas taip pat rodo, kad net ir skeptiškai nusiteikę respondentai pripažino tam tikrą šaltinio ir turinio atitikimo lygį, o pati skalė fiksuoja labiau vidutinį nei krašutinį požiūrį.

Siekiant patikrinti tyrimo kintamųjų pasiskirstymo normalumą, buvo atlikti Kolmogorovo–Smirnov ir Shapiro–Wilko testai (žr. 15 lentelė). Šie testai leidžia įvertinti, ar duomenų skirstinys reikšmingai skiriasi nuo teorinio normaliojo pasiskirstymo, o tai svarbu planuojant tolesnės analizės metodus.

15 lentelė. Tyrimo kintamųjų pasiskirstymo analizė naudojantis Kolmogorovo–Smirnov ir Shapiro–Wilko testais (N = 372)

Skalės pavadinimas	Dimensijos	Kolmogorovo–Smirnov testas			Shapiro–Wilko testas		
		Statistika	Laisvės laipsniai (df)	p reikšmė	Statistika	Laisvės laipsniai (df)	p reikšmė
Šaltinio ir turinio atitikimas		0,101	372	<0,001	0,974	372	<0,001
Turinio patikimumas		0,109	372	<0,001	0,965	372	<0,001
Įsitraukimo ketinimai		0,116	372	<0,001	0,926	372	<0,001
	Šaltinio kompetencija (6)	0,072	372	<0,001	0,970	372	<0,001

Šaltinio patikimumas	Šaltinio palankumas (3)	0,152	372	<0,001	0,966	372	<0,001
	Šaltinio empatija (3)	0,125	372	<0,001	0,970	372	<0,001
	Pasitikėjimas šaltiniu (3)	0,125	372	<0,001	0,956	372	<0,001

Remiantis rezultatais, visų tirtų konstrukty – šaltinio ir turinio atitikimo, turinio patikimumo, ištraukimo ketinimų, šaltinio kompetencijos, palankumo, empatijos ir pasitikėjimo – reikšmės neatitiko normalumo kriterijų. Tai patvirtina abiejų testų rezultatai, kurių p reikšmės buvo mažesnės nei 0,001, o tai rodo statistiškai reikšmingą nukrypimą nuo normaliojo pasiskirstymo.

Tyrimo duomenims toliau analizuoti buvo atlikta Spearmano rangų koreliacijos analizė (žr. 16 lentelė ir 18 priedas), kurios tikslas – įvertinti statistinių ryšių tarp pagrindinių tiriamųjų konstrukty stiprumą (Pallant, 2020). Kadangi ankstesnė normalumo patikra (Kolmogorovo–Smirnov ir Shapiro–Wilko testai) parodė, jog visų konstrukty duomenys reikšmingai skiriasi nuo normaliojo pasiskirstymo ($p < 0,001$), pasirinktas neparametrinis metodas. Šio tyrimo koreliacinių sąsajų interpretacijai pasitelkta konceptuali reikšmių klasifikacija, paremta Cohen'o et al. (2003) rekomendacijomis, kurias adaptavo Piligrimienė (2016). Pagal šį principą, sąryšis tarp dviejų kintamųjų laikomas labai silpnu, kai koreliacijos koeficientas neviršija 0,19, silpnu – kai svyruoja tarp 0,20 ir 0,39, vidutinio stiprumo – nuo 0,40 iki 0,69, stipriu – tarp 0,70 ir 0,89, o labai stipriu – kai reikšmė viršija 0,90. Ši klasifikacija leido sistemingai įvertinti ryšių stiprumą tarp analizuojamų konstrukty. Dvi žvaigždutės (**) prie reikšmių nurodo, kad šie rezultatai yra statistiškai reikšmingi ($p < 0,001$).

16 lentelė. Kintamųjų koreliacijos analizė (N = 372)

Kintamasis	Šaltinio ir turinio atitikimas	Turinio patikimumas	Ištraukimo ketinimai	Šaltinio kompetencija	Šaltinio palankumas	Šaltinio empatija	Pasitikėjimas šaltiniu
Šaltinio ir turinio atitikimas	1,000 (labai stipri)	0,457** (vidutinė)	0,244** (silpna)	0,405** (vidutinė)	0,292** (silpna)	0,436** (vidutinė)	0,364** (silpna)
Turinio patikimumas		1,000 (labai stipri)	0,374** (silpna)	0,566** (vidutinė)	0,356** (silpna)	0,599** (vidutinė)	0,563** (vidutinė)
Ištraukimo ketinimai			1,000 (labai stipri)	0,283** (silpna)	0,400** (vidutinė)	0,455** (vidutinė)	0,221** (silpna)
Šaltinio kompetencija				1,000 (labai stipri)	0,449** (silpna)	0,555** (vidutinė)	0,629** (vidutinė)
Šaltinio palankumas					1,000 (labai stipri)	0,427** (silpna)	0,455** (vidutinė)
Šaltinio empatija						1,000 (labai stipri)	0,570** (vidutinė)
Pasitikėjimas šaltiniu							1,000 (labai stipri)

Tyrimo metu nustatyta, kad visi kintamieji koreliuoja tarpusavyje teigiama kryptimi – tai rodo, kad aukštesnis vieno konstrukto įvertinimas yra susijęs su aukštesniais kitų kintamųjų vertinimais. Statistiškai reikšmingi ryšiai aptikti beveik tarp visų tiriamųjų porų, o jų stiprumas dažniausiai pasiekė vidutinio ar silpno lygmens ribas. Stipriausias sąryšis identifikuotas tarp šaltinio kompetencijos ir pasitikėjimo juo (0,629**), taip pat tarp turinio patikimumo ir šaltinio empatijos (0,599**), bei tarp turinio patikimumo ir šaltinio pasitikėjimo (0,563**). Tai leidžia daryti prielaidą,

kad emocinės savybės, tokios kaip šaltinio jautrumas ir supratingumas, glaudžiai siejasi su bendru pasitikėjimo šaltiniu lygiu. Taip pat pastebimas tvirtas ryšys tarp šaltinio empatijos ir pasitikėjimo ($0,570^{**}$), o tai rodo, kad emocinis šaltinio aspektas, susijęs su gebėjimu atjausti auditoriją, tiesiogiai veikia jo patikimumo suvokimą.

Vidutinio stiprumo koreliacijos taip pat nustatytos tarp šaltinio ir turinio atitikimo bei šaltinio kompetencijos ($0,405^{**}$), tarp šaltinio ir turinio atitikimo bei šaltinio empatijos ($0,436^{**}$), taip pat tarp šaltinio kompetencijos ir empatijos ($0,555^{**}$), ir šaltinio palankumo bei pasitikėjimo ($0,455^{**}$). Šie ryšiai rodo, kad vartotojai vertina šaltinį kaip kompetentingą ir patikimą tuomet, kai jo publikuojamas turinys yra nuoseklus, turi sąsajų su šaltiniu.

Silpno stiprumo, tačiau reikšmingi ryšiai stebimi tarp šaltinio ir turinio atitikimo bei turinio patikimumo ($0,457^{**}$), tarp šaltinio ir turinio atitikimo bei šaltinio pasitikėjimo ($0,364^{**}$), tarp turinio patikimumo ir šaltinio palankumo ($0,356^{**}$), taip pat tarp įsitraukimo ketinimų ir šaltinio palankumo ($0,400^{**}$) ar pasitikėjimo ($0,221^{**}$). Tokie rezultatai liudija, kad nors kai kurios sąsajos tarp konstrukto nėra itin stiprios, jos vis tiek rodo prasmingą ryšį – pavyzdžiui, kai respondentai mato šaltinį kaip nuoseklų ir patikimą, jie dažniau nori įsitraukti į jo siūlomą veiklą ar turinį.

Mažiausio stiprumo, bet vis dar teigiama koreliacija stebima tarp šaltinio ir turinio atitikimo bei įsitraukimo ketinimų ($0,244^{**}$). Šis rezultatas leidžia manyti, kad nors šaltinio ir turinio atitikimas gali turėti įtakos vartotojo ketinimui reaguoti ar įsitraukti į turinį, šis poveikis nėra dominuojantis – labiau priklauso nuo sudėtingesnės, galbūt ir individualizuotos, sąveikos su šaltiniu. Visi šie rezultatai patvirtina ankstesnėje faktorinėje analizėje išskirtų konstrukcijų empirinį pagrįstumą ir tarpusavio sąryšius.

Siekiant nustatyti, ar respondentų lytis, amžius, išsilavinimas ar pajamos daro įtaką jų vertinimams, susijusiems su vartotojų atsaku į DI sukurtą turinį socialiniuose tinkluose, buvo atlikti neparametriniai Kruskal–Wallis H testai (žr. 18 priedas). Kadangi ankstesniais testais nustatyta, jog konstrukto reikšmės neatitinka normaliojo pasiskirstymo prielaidų, palyginimams tarp grupių buvo pasirinktas būtent šis metodas, leidžiantis įvertinti daugiau nei dviejų nepriklausomų grupių vertinimų skirtumus. Buvo analizuojama, ar respondentai pagal lytį, amžių, išsilavinimą ir gaunamas pajamas reikšmingai skiriasi vertindamos tokius konstrukto aspektus kaip šaltinio ir turinio atitikimas, turinio patikimumas, įsitraukimo ketinimai bei šaltinio patikimumo dimensijos – kompetencija, palankumas, empatija ir pasitikėjimas. Gauti rezultatai rodo, kad statistiškai reikšmingas skirtumas pagal lytį buvo nustatytas tik viename konstrukto aspekte – šaltinio empatijos skalėje, kurioje p reikšmė buvo mažesnė už $0,05$ ($p=0,016$). Tai reiškia, kad respondentai reikšmingai skyrėsi savo vertinimuose priklausomai nuo lyties, kai buvo vertinamas rūpestingumas vartotojo ir jo interesų atžvilgiu. Nors reikšmingų skirtumų nenustatyta likusiuose konstrukto aspektuose, pastebimos aiškos tendencijos: moterų vertinimai nuosekliai viršija vyrų beveik visose skalėse, ypač šaltinio palankumo, pasitikėjimo ir kompetencijos dimensijose. Gauti rezultatai tai pat parodė, kad statistiškai reikšmingi skirtumai tarp amžiaus grupių buvo nustatyti beveik visose konstrukto srityse, išskyrus šaltinio pasitikėjimą ($p=0,236$). Reikšmingiausias skirtumas užfiksuotas įsitraukimo ketinimų skalėje ($p<0,001$), kurioje vyresni respondentai (32–67 m.) turėjo aukščiausią rangų vidurkį (243,63), o jauniausi (15–23 m.) – žemiausią (151,38). Tai rodo, kad vyresni asmenys dažniau teigė ketinantys reaguoti į DI sukurtą turinį (dalintis, komentuoti, spausti „patinka“), nei jaunesni. Taip pat reikšmingi skirtumai užfiksuoti turinio patikimumo ($p<0,001$), šaltinio ir turinio atitikimo ($p=0,008$), šaltinio empatijos ($p=0,007$), kompetencijos ($p=0,033$) ir palankumo ($p=0,039$) skalėse. Visose šiose

dimensijose pastebima tendencija, kad vyresni respondentai (ypač grupė 32–67 m.) nuosekliai priskiria aukštesnius vertinimus, nei jaunesni. Remiantis respondentų vertinimų analize pagal išsilavinimą, pastebėti statistiškai reikšmingi skirtumai tarp skirtingų išsilavinimo grupių nustatyti dvejose konstrukto skalėse: vartotojų įsitraukimo ketinimų ($p=0,001$) ir šaltinio palankumo ($p=0,006$). Tai leidžia teigti, kad respondentų išsilavinimas gali turėti reikšmingos įtakos tiek jų elgesio ketinimams reaguoti į DI turinį, tiek vertinimui, ar šaltinis yra supratingas ir jautrus. Abu atvejai išsiskyrė tuo, kad žemesnio išsilavinimo respondentai (ypač su pagrindiniu ar aukštesniu išsilavinimu) turėjo aukštesnes rangų reikšmes nei aukštąjį išsilavinimą įgiję dalyviai. Be to, statistiškai reikšmingi skirtumai tarp pajamų grupių buvo nustatyti penkiose iš septynių vertintų konstrukto dimensijų: šaltinio ir turinio atitikimo ($p=0,009$), šaltinio kompetencijos ($p=0,003$), šaltinio palankumo ($p=0,019$) ir šaltinio pasitikėjimo ($p=0,033$) kintamuosiuose. Įdomu tai, kad kai kuriose skalėse aukštesni pajamų lygiai nebūtinai reiškė aukštesnius šaltinio ar turinio vertinimus – priešingai, aukščiausias pajamas nurodę respondentai dažnai priskirdavo žemesnius įvertinimus, ypač šaltinio palankumo ir atitikimo dimensijose. Tai gali reikšti kritiškesnę ar atsargesnę požiūrį į DI turinį tarp didesnes pajamas gaunančių vartotojų.

Apibendrinant galima teigti, kad sociodemografiniai veiksniai – lytis, amžius ir pajamos – darė reikšmingą įtaką tam tikriems respondentų vertinimo aspektams, susijusiems su DI sukurtu turiniu ir jo šaltiniu. Nors ne visose konstrukto dimensijose skirtumai buvo statistiškai reikšmingi, daugeliu atvejų stebėtos tendencijos leidžia identifikuoti ryškius skirtumus tarp grupių. Lyties analizė parodė, kad moterys nuosekliai skyrė aukštesnius vertinimus, ypač šaltinio empatijos skalėje, o tai leidžia daryti prielaidą, kad emocinis jautrumas ir pasitikėjimas šaltiniu gali būti glaudžiau susijęs su moterų komunikacinėmis nuostatomis. Amžiaus grupių palyginimas atskleidė aiškią tendenciją: vyresni respondentai dažniau teigiamai vertino turinį, jo patikimumą bei buvo linkę išreikšti didesnę įsitraukimo tikimybę, o taip pat labiau linko pozityviai vertinti šaltinio emocines savybes. Šie rezultatai leidžia manyti, kad amžius yra svarbus veiksnys, formuojantis ne tik pasitikėjimą DI komunikacija, bet ir norą aktyviai į ją įsitraukti. Panašūs skirtumai buvo stebėti ir pagal pajamų lygį. Nors rezultatai čia buvo kiek nevienareikšmiai, ypač aiškūs buvo skirtumai šaltinio kompetencijos, pasitikėjimo ir palankumo skalėse. Pastebėta, kad žemesnes pajamas gaunantys asmenys dažniau šaltinį vertino palankiau ir kompetentingiau nei aukštesnes pajamas deklaravę respondentai, o tai gali rodyti skirtingą informacijos vertinimo jautrumą, susijusį su ekonominiu kontekstu. Bendros analizės rezultatai leidžia daryti išvadą, kad respondentų socialiniai ir demografiniai bruožai veikia jų reakcijas į DI generuotą turinį, ypač tada, kai vertinami pasitikėjimo, empatijos ir įsitraukimo aspektai. Tai patvirtina, kad DI komunikacijos efektyvumas gali priklausyti ne tik nuo paties turinio, bet ir nuo auditorijos savybių, todėl sociodemografinės charakteristikos turėtų būti laikomos reikšmingu kontekstiniu veiksnio šiuolaikinėje komunikacijos analizėje.

4.3. Eksperimento grupių homogeniškumo ir manipuliacijų veiksmingumo analizė

Prieš hipotezių tinkrinimą, tikslinga atlikti manipuliacijų veiksmingumo bei eksperimento grupių homogeniškumo analizes, siekiant užtikrinti tyrimo sąlygų patikimumą ir rezultatų pagrįstumą.

Eksperimento grupių homogeniškumas. Siekiant įsitikinti, kad eksperimentinių grupių, kurioms buvo pateikti skirtingi scenarijai, sudėtis yra homogeniška pagal pagrindinius demografinius požymius, buvo atlikta respondentų pasiskirstymo analizė pagal lytį, amžių, pajamas ir išsilavinimą. Duomenų homogeninio pasiskirstymo patikrinimas buvo atliktas pirmausia įvertinant lyties kintamąjį. Kadangi tik du respondentai iš 372 pasirinko atsakymo variantą „kita“, jie buvo pašalinti iš analizės, siekiant

užtikrinti patikimą statistinį įvertinimą (N = 370). Pasiskirstymo pagal lytį homogeniškumui įvertinti buvo taikytas Pearsono chi kvadrato testas (žr. 17 lentelė), leidžiantis nustatyti, ar tarp dviejų kintamųjų egzistuoja statistiškai reikšmingas ryšys (Pallant, 2020), bei dažnių analizė (angl. *crosstabulation*).

17 lentelė. Pearsono chi kvadrato testo rezultatai analizuojant respondentų pasiskirstymo eksperimentinėse grupėse homogeniškumą (N = 370)

Pearsono chi kvadrato testas	Reikšmė	Laisvės laipsniai (df)	p reikšmė
Pearsono chi kvadratas	0,190a	3	0,979
Tikimybių santykio testas	0,190	3	0,979
Linijinės priklausomybės testas	0,175	1	0,676

Chi kvadrato testo rezultatai parodė, kad p reikšmė yra didesnė nei 0,05, o tai rodo, jog nėra reikšmingų skirtumų tarp respondentų lyties ir grupės priskyrimo. Tai rodo, kad tiriamieji eksperimento sąlygoms buvo priskirti tolygiai pagal lyties požymį. Analizuojant procentinį pasiskirstymą matyti, kad visose keturiose eksperimentinėse grupėse vyrų ir moterų santykiai buvo labai panašūs. Pavyzdžiui, grupėje „su žymėjimu, atitinka DI“ moterų sudarė 58,3 %, o vyrų – 41,7 %, o grupėje „be žymėjimo, atitinka DI“ atitinkamai 58,8 % ir 41,2 %. Likusiose grupėse skirtumai taip pat nedideli (pavyzdžiui, „su žymėjimu, neatitinka DI“ – 61,5 % moterų ir 38,5 % vyrų; „be žymėjimo, neatitinka DI“ – 59,5 % moterų ir 40,5 % vyrų). Tokie rezultatai patvirtina, kad eksperimentinių sąlygų grupės buvo sudarytos atsitiktinai ir vienodai pagal tiriamųjų lytį, o tai užtikrina tyrimo patikimumą ir galimybę objektyviai lyginti skirtingų grupių reakcijas.

Siekiant įvertinti, ar tyrimo dalyviai keturiose eksperimentinėse grupėse pagal amžių yra pasiskirstę tolygiai, buvo atlikta dispersijos analizė (ANOVA) (žr. 19 priedas). Prieš tai apskaičiuoti amžiaus aprašomosios statistikos rodikliai leido nustatyti, kad vidutinis respondentų amžius grupėse svyravo nuo 28,40 iki 31,08 metų. Didžiausias vidurkis buvo grupėje, kurioje pateiktas pažymėtas įrašas apie su DI susijusį produktą – jų amžiaus vidurkis siekė 31,08 metų. Grupėje, kurioje pateiktas pažymėtas įrašas apie nesusijusį su DI turinį, vidurkis buvo 28,81 metų. Nepažymėto įrašo apie DI inovaciją grupės vidutinis amžius buvo 29,35 metų, o mažiausias vidurkis užfiksuotas nepažymėto įrašo apie su DI nesusijusį produktą grupėje – 28,40 metų. Apibendrinus visus respondentus, bendras amžiaus vidurkis siekė 29,44 metų (SD = 10,08).

Prieš atliekant dispersijos analizę (ANOVA), buvo patikrinta lygių dispersijų prielaida naudojant Levene'o testą (žr. 18 lentelė). Visi rezultatai buvo statistiškai nereikšmingi ($p > 0,05$), todėl galima teigti, kad dispersijos tarp grupių nesiskyrė. ANOVA rezultatai parodė, kad amžiaus vidurkiai tarp eksperimentinių grupių statistiškai reikšmingai nesiskyrė ($F(3, 368) = 1,341, p=0,261$). Tai reiškia, kad respondentų grupės buvo homogeniškos pagal amžiaus požymį.

18 lentelė. Lygių dispersijų prielaidos Lavene'o testas: amžius (N = 372)

Variacijų homogeniškumo testai		Levene'o statistika	Laisvės laipsniai (df1)	Laisvės laipsniai (df2)	p reikšmė (Sig.)
Koks yra jūsų amžius?	Pagal vidurkį	1,738	3	368	0,159
	Pagal medianą	0,764	3	368	0,515
	Pagal medianą (pakoreguoti df)	0,764	3	354,255	0,515

	Pagal nukirptą vidurkį	1,523	3	368	0,208
--	------------------------	-------	---	-----	-------

Vertinant, ar eksperimentinėse grupėse respondentai pagal išsilavinimą buvo pasiskirstę tolygiai, buvo atliktas neparametrinis Kruskal–Wallis testas. Analizės rezultatai parodė, kad vidurkių rangų (žr. 19 lentelė) skirtumai tarp grupių buvo nežymūs – aukščiausias rangas (195,62) priklausė respondentams, kurie matė pažymėtą įrašą apie DI inovaciją, o žemiausias (178,34) – tiems, kurie vertino pažymėtą, bet su DI nesusijusį turinį. Vis dėlto Kruskal–Wallis testo rezultatas ($H = 1,454$, $p=0,693$) rodo, kad šie skirtumai nėra statistiškai reikšmingi. Kadangi p reikšmė yra didesnė nei 0,05, galima daryti išvadą, jog eksperimentinės grupės buvo homogeniškos pagal išsilavinimą.

19 lentelė. Išsilavinimo rangų vidurkiai pagal eksperimentines grupes

Vertinimo kriterijus	Eksperimentinė grupė	N	Vidurkių rangai
Koks yra aukščiausias jūsų įgytas išsilavinimas?	Su žymėjimu, atitinka DI	98	195,62
	Be žymėjimo, atitinka DI	102	183,23
	Su žymėjimu, neatitinka DI	64	178,34
	Be žymėjimo, neatitinka DI	108	186,15

Siekiant įvertinti, ar respondentai yra pasiskirstę tolygiai eksperimentinėse grupėse pagal jų mėnesinės pajamas, buvo atlikta kryžminių dažnių analizė (angl. *Crosstabulation*) (žr. 19 priedas). Rezultatai parodė, kad visose keturiose grupėse dominuoja respondentai, uždirbantys nuo 1000 iki 1499 eurų (bendras dažnis – 30,6 %) bei nuo 1500 iki 1999 eurų (18,3 %). Žemiausio pajamų lygmens – mažiau nei 500 eurų – respondentai sudarė tik 7 % visų dalyvių, tuo tarpu aukščiausio lygmens – 2500 eurų ar daugiau – 8,1 %. Taip pat 15,6 % respondentų atsisakė nurodyti savo pajamas. Procentinė analizė rodo, kad grupių sudėtis pagal pajamas yra gana panaši. Šie rezultatai leidžia daryti išvadą, kad respondentai pagal pajamas tarp eksperimentinių grupių buvo pasiskirstę palyginti tolygiai. Tai rodo, kad eksperimento grupių homogeniškumas pagal šį demografinį požymį yra pakankamas.

Manipuliacijos efektyvumas. Siekiant patikrinti manipuliacijos efektyvumą – įvertinti, ar skirtingai pateiktas turinys (su DI autorystės žymėjimu ar be jo) ir turinio tematika (susijusi su DI inovacija ar nesusijusi) paveikė respondentų įsitikinimą, kad įrašą sukūrė DI – buvo atlikta kryžminė dažnių analizė ir Pearsono chi kvadrato testas. Analizėje buvo lyginami du kintamieji: ar įrašas buvo pažymėtas kaip sukurtas DI (taip arba ne), ir kaip respondentai įvertino įrašo autorystę (taip, ne arba nežinau). Gauti rezultatai atskleidė stiprią sąsają tarp šių kintamųjų.

Pearsono chi kvadrato testas (žr. 20 lentelė) parodė, kad tarp kintamųjų egzistuoja statistiškai reikšmingas ryšys ($\chi^2(2) = 155,916$, $p<0,001$). Tai leidžia teigti, kad įrašo žyma kaip sukurtas DI reikšmingai padidino tikimybę, kad respondentai šio įrašo autorystę priskirs DI. Šį rezultatą sustiprino papildomi testai: tikimybių santykio testo (angl. *likelihood ratio*) reikšmė siekė 214,339 ($p<0,001$), o linijinės priklausomybės testo – 18,298, ($p<0,001$). Šie du rodikliai dar labiau patvirtina statistinį ryšį tarp eksperimentinio manipuliacinio elemento ir respondentų vertinimų.

20 lentelė. Pearsono chi kvadrato testo rezultatai analizuojant turinio sąsąjas su DI ir respondentų autorystės vertinimą (N = 372)

Pearsono chi kvadrato testas	Reikšmė	Laisvės laipsniai (df)	p reikšmė
Pearsono chi kvadratas	155,916 ^a	2	<0,001

Tikimybių santykio testas	214,339	2	<0,001
Linijinės priklausomybės testas	18,298	1	<0,001

Kryžminė dažnių analizė parodė reikšmingą skirtumą tarp grupių (žr. 21 lentelė). Tarp tų respondentų, kurie matė pažymėtą įrašą, net 39,5 proc. nurodė, kad įrašą sukūrė DI, o 60,5 proc. pažymėjo, kad nežino. Priešingai, tarp tų, kurie matė nepažymėtą įrašą, net 44,8 proc. nurodė, kad įrašo nesukūrė DI, o 55,2 proc. pasirinko atsakymą „nežinau“. Toks pasiskirstymas rodo, kad įrašo žymėjimas etikete „AI info“ reikšmingai sustiprina suvokimą apie DI autorystę. Be to, kryžminės lentelės duomenys atskleidė, kad bendras neapsisprendusių (atsakiusių „nežinau“) dalyvių skaičius siekė 57,5 proc. visų respondentų. Didelis respondentų, pasirinkusių atsakymą „nežinau“, skaičius rodo, kad vartotojai dažnai nesijautė pakankamai užtikrinti priimdami sprendimą dėl įrašo autorystės. Viena iš galimų tokio rezultato priežasčių yra ta, kad manipuliacijos klausimas, kuriuo buvo tikrinama, ar respondentai pastebėjo „AI info“ žymą, buvo pateiktas apklausos pabaigoje. Todėl tikėtina, kad dalis dalyvių galėjo nepastebėti žymėjimo, kuris, lyginant su įrašo tekstu, buvo mažesnis arba tiesiog jį pamiršti.

21 lentelė. Kryžminės dažnių analizės rezultatai vertinant turinio sąsajos su DI ir respondentų autorystės sąveiką (N = 372)

	Respondentų atsakymai į klausimą: Įrašas buvo pažymėtas dirbtinio intelekto autorystę nurodančia etikete „AI info“ (liet. dirbtinio intelekto informacija).			
DI žyma	Taip	Ne	Nežinau	Iš viso
Yra	64 (39,5 %)	0 (0,0 %)	98 (60,5 %)	162 (100 %)
Nėra	0 (0,0 %)	94 (44,8 %)	116 (55,2 %)	210 (100 %)
Iš viso	64 (17,2 %)	94 (25,3 %)	214 (57,5 %)	372 (100 %)

Vertinant bendrąjį atsakymų pasiskirstymą (žr. 22 lentelė), respondentų (N = 372) vidutinis vertinimas, kad socialinių tinklų įrašą sukūrė DI, siekė 4,73 (SD = 1,610). Detaliau analizuojant pagal grupes, pastebėta, kad aukščiausias vertinimo vidurkis buvo tarp respondentų, kurie matė su žymėjimu pateiktą įrašą apie nesusijusį su DI produktą (M = 4,94, SD = 1,531). Antroje vietoje pagal vidurkį buvo tie, kurie matė be žymėjimo pateiktą įrašą apie nesusijusį produktą (M = 4,87, SD = 1,773), o po jų – respondentai, kurie matė su žymėjimu pateiktą įrašą apie DI inovaciją (M = 4,86, SD = 1,450). Mažiausias vertinimo vidurkis užfiksuotas tarp tų, kurie matė be žymėjimo pateiktą įrašą apie DI inovaciją (M = 4,31, SD = 1,573).

22 lentelė. Respondentų vertinimų pasiskirstymas dėl DI įrašo autorystės pagal eksperimentines grupes (N = 372)

Eksperimentinė grupė	Vertinimo vidurkis	Dažnis (N)	Standartinis nuokrypis (SD)
Su žymėjimu, atitinka DI	4,86	98	1,450
Be žymėjimo, atitinka DI	4,31	102	1,573
Su žymėjimu, neatitinka DI	4,94	64	1,531
Be žymėjimo, neatitinka DI	4,87	108	1,773

Siekiant įvertinti, ar respondentų vertinimo, kad socialinių tinklų įrašą sukūrė DI, vidurkiai statistiškai reikšmingai skiriasi tarp keturių eksperimentinių grupių, buvo atlikta vienfaktorinė dispersijos analizė (ANOVA) (žr. 20 priedas). ANOVA rezultatai parodė, kad grupių vidurkiai skiriasi statistiškai reikšmingai, $F(3,368) = 3,157$, $p=0,025$. Tai leidžia daryti prielaidą, jog bent viena grupė skyrėsi nuo kitų pagal pritarimo lygį. Efekto dydžio analizė parodė, kad $\eta^2 = 0,025$, kas rodo mažo stiprumo efektą (pagal Cohen'o (1988) rekomendacijas, mažas efektas laikomas nuo 0,01 iki 0,06).

Apibendrinant vidurkių analizę, matyti, kad didžiausias vertinimas apie įrašo autorystės priskyrimą DI buvo tarp respondentų, kurie matė nepažymėtą įrašą apie nesusijusį su DI produktą ($M = 4,94$, $SD = 1,531$), o mažiausias – tarp tų, kurie vertino be žymėjimo pateiktą įrašą apie DI inovaciją ($M = 4,31$, $SD = 1,573$). Kitose grupėse vertinimų vidurkiai siekė atitinkamai 4,86 ir 4,87 balus. Rezultatai rodo, kad aiškus turinio žymėjimas DI autorystės etikete sistemingai stiprina vartotojų suvokimą apie DI sukurtą turinį, o žymėjimo nebuvimas, net atvejais, kai turinio tematika tiesiogiai susijus su DI, nesukuria stipresnio vartotojų įsitikinimo apie DI autorystę. Tokie rezultatai leidžia daryti išvadą, kad vartotojai labiau remiasi pateiktomis aiškiomis nuorodomis apie turinio kilmę nei savarankiškai priskiria autorystę remdamiesi tik turinio tematika.

4.4. Atvirų atsakymų, nurodančių vartotojų emocinį atsaką į manipuliacinį turinį, sentimentų analizė

Siekiant papildomai įvertinti vartotojų emocinį atsaką į pateiktus socialinių tinklų įrašų scenarijus, tyrime buvo pasirinkta atlikti atvirų atsakymų sentimentų analizę. Šią analizę leido įgyvendinti anketos manipuliacijos patikros klausimas, kuriame respondentų buvo prašoma pagrįsti savo nuomonę apie įrašo autorystę. Kadangi šis klausimas skatino dalyvius pateikti savarankišką vertinimą, surinkti atsakymai sudarė tinkamą bazę emocinio tono ir vartotojų požiūrio kryptingumui analizuoti.

Pradiniame duomenų rinkinyje buvo užfiksuoti 372 atviri atsakymai. Siekiant užtikrinti analizės kokybę ir pašalinti nereikšmingą informaciją, buvo atliktas duomenų filtravimas. Iš tolesnės analizės buvo pašalinti atsakymai, kuriuose vietoje nuomonės buvo pateikti tik emocijų piktogramos (emotikonai), pavieniai simboliai, taškai ar pavienės raidės. Tokie įrašai buvo laikomi neturinčiais analitinės vertės vartotojų požiūrio interpretacijai. Atlikus šį filtravimą, galutinei sentimentų analizei buvo atrinkti 180 pilnaverčių tekstinių atsakymų (žr. 21 priedas).

Siekiant giliau suprasti vartotojų vertinimo priežastis ir emocinį atsaką į pateiktus socialinių tinklų įrašų scenarijus, buvo atlikta tekstinių atvirų atsakymų analizė, kuri buvo organizuota trimis etapais. Pirmiausia, analizė buvo atlikta pasitelkiant vizualinės analizės įrankį *Wordclouds.com*. Į šią platformą buvo įkelti iš anksto apdoroti atviri respondentų atsakymai, pašalinus nereikšmingus elementus, tokius kaip vienos raidės žymėjimai, simboliai bei jungtukai. Iš viso analizei buvo pateikti 180 atvirų atsakymų, sudarytų iš 954 unikalių žodžių. Tekstinių duomenų analizės metu žodžiai buvo sujungti į bendrines formas, pašalinant linksniuotės skirtumus. Šis žingsnis buvo atliktas siekiant užtikrinti tikslesnį dažniausiai vartojamų sąvokų identifikavimą ir sumažinti reikšmių išskaidymą dėl lietuvių kalbai būdingos morfologinės įvairovės. Tokiu būdu, skirtingos to paties žodžio formos (pavyzdžiui, „tekstas“, „teksto“, „tekste“) buvo apjungtos į vieną bazinę formą, leidžiančią tiksliau atspindėti nagrinėjamų temų pasikartojimą ir ryškumą žodžių debesyje.

Analizuojant sugeneruotą žodžių debesį (žr. 4 pav.), matyti, kad dominuojančios sąvokos yra tiesiogiai susijusios su teksto turiniu, DI technologijomis bei vertinimo aspektais. Dažniausiai

Taip pat pastebėta, kad tarp dažniausiai minimų žodžių atsiranda tokie elementai kaip „vaizdas“, „skamba“, „klaidos“ ir net specifinė žyma „AI info“. Šie elementai rodo, kad dalyviai atkreipė dėmesį ne tik į teksto kokybę, bet ir į papildomas įrašo sudėtinės dalis, kurios galėjo lemti jų suvokimą apie turinio autorystę.



4 pav. Semantinės respondentų atvirų atsakymų analizės žodžių debesis (N = 180)

Antruoju etapu buvo atlikta autorystės priskyrimo ir emocinio tono analizė (žr. 23 lentelė). Atsakymai buvo vertinami pagal jų turinį ir semantinę raišką, siekiant nustatyti emocinį atsako toną: teigiamą, neutralų arba neigiamą:

- neigiama reakcija buvo laikoma tokia, kai dalyvis pabrėždavo turinio stilistinius, gramatinius ar natūralumo trūkumus, aiškiai rodydamas kritinį požiūrį,
- neutralūs atsakymai buvo tie, kurie pateikė pastebėjimus be aiškos emocinės konotacijos,
- teigiami atsakymai – tie, kurie išreiškė pasitenkinimą įrašo kokybe ar natūralumu.

23 lentelė. Atvirų respondentų atsakymų autorystės priskyrimo ir emocinio tono analizė (N = 180)

Vertinimo dimensija	Kategorija	Dažnis (N)	Procentinė dalis (proc.)
Autorystės priskyrimas	Dirbtinis intelektas	115	63,9 proc.
	Žmogus	18	10,0 proc.
	Abu (žmogus ir DI)	20	11,1 proc.
	Neaišku	27	15,0 proc.
Emocinis tonas	Teigiamas	10	5,6 proc.
	Neutralus	122	67,8 proc.
	Neigiamas	48	26,7 proc.

Autorystės vertinimo rezultatai parodė, kad dauguma respondentų įrašus priskyrė DI autorystei. 63,9 proc. (N = 115) atsakymų nurodė DI kaip galimą autorių. Žmogaus autorystė buvo įvardyta tik 10 proc. (N = 18) atsakymų, o 11,1 proc. (N = 20) respondentų nurodė, kad įrašą ar jo dalis galėjo kurti žmogus kartu su DI pagalba. Likusieji 15 proc. (N = 27) atsakymų nepateikė aiškos nuomonės dėl įrašo kilmės – arba manė, jog įrašą galėjo kurti tiek žmogus, tiek DI, arba nepateikė su įrašo autoryste susijusios nuomonės.

Emocinio tono analizė atskleidė, kad didžioji dalis atsakymų (67,8 proc., N = 122) buvo neutralūs, t. y., vartotojai pateikė pastebėjimus be aiškos emocinės konotacijos. Neutralūs atsakymai buvo orientuoti į faktinį turinio vertinimą be stiprių emocijų. Šie atsakymai dažnai atspindėjo neapsisprendimą dėl įrašo autorystės arba subalansuotą požiūrį, kai buvo išvardijami tiek žmogui, tiek DI būdingi bruožai. Vienas respondentas rašė: „*Sunku atskirti ar sukurta žmogaus, ar dirbtinio intelekto.*“, kitas pažymėjo: „*Gali būti ir žmogus, ir AI, visi tie marketinginiai tekstai panašūs.*“ Pastebėta, kad neutralūs vertinimai dažnai susiję su pastebimais teksto stilistikos ar rašybos požymiais, pavyzdžiui: „*Emoji ir brūkšnių naudojimas bei pats išdėstymas tekste kelia įtarimą, nes taip daro dažnai dirbtinis intelektas.*“ Ši respondentų grupė nesirėmė stipria emocine reakcija, o analizavo turinį techniniu ir lingvistiniu aspektu.

Neigiamą emocinį toną – dažniausiai susijusį su gramatikos, stilistikos klaidomis ar natūralumo trūkumų akcentavimu – išreiškė 26,7 proc. (N = 48) respondentų. Respondentai pastebėjo, kad tiek tekstas, tiek vizualinė medžiaga kartais atrodė nenatūraliai ar neprofesionaliai. Kaip pažymėjo vienas dalyvis, „*Laužyta ir nerišli kalba*“, o kitas pabrėžė, kad „*Tekstas neskamba 'žmogiškai', atrodo toks... perdėtai išstbulintas, nuotrauka irgi neįtikina.*“ Kritiškas vertinimas buvo siejamas ir su loginėmis spragomis: „*Pirmajame sakinyje galima pastebėti nelogišką priežasties-pasekmės sąsają: ar tikrai KTU sukūrė įrenginį, nes depresija paliečia milijonus žmonių.*“ Neigiamos reakcijos taip pat buvo nukreiptos į nuotraukos kokybę ir atitikimą realybei: „*Nuotrauka atrodo labai redaguota ir pats tekstas neskamba tikroviškai ar lyg būtų rašytas žmogaus.*“

Tik 5,6 proc. (N = 10) atsakymų pasižymėjo teigiamu tonu, nurodant teksto aiškumą, sklandumą ir bendrą profesionalią raišką. Respondentai, kurie pateikė teigiamus atsakymus, dažniausiai pabrėždavo, kad tekstas skaitomas lengvai ir be klaidų. Kaip rašė vienas dalyvis, „*Tekstas sklandus, aiškus, išsamus. Nenaudojami hiperbolizmai ar nerodoma per daug entuziazmo, kuriuo pasižymi DI.*“ Kitas respondentų pastebėjimas akcentavo teksto vientisumą ir struktūrą: „*Aprašyta be klaidų, logiškai.*“ Teigiami vertinimai buvo išreiškiami ir dėl institucijos patikimumo: „*Prie įrašo nėra nurodyta, kad ji sukūrė dirbtinis intelektas. Matydama, kad įrašas yra iš mokslo įstaigos – laikau įrašą patikimesniu ir manau, kad komunikuoja universiteto atstovai.*“ Šie atsakymai leidžia daryti išvadą, kad profesionaliai pateiktas turinys, net ir galimai sukurtas DI, sukelia teigiamą vartotojų reakciją.

Analizės metu paaiškėjo ir papildoma svarbi įžvalga, susijusi su vartotojų lūkesčiais institucijoms dėl skaidrumo komunikacijoje. Vienas iš respondentų nurodė: „*Nuotrauka atrodo, lyg galėtų būt padaryta AI, bet iš KTU tikėčiausi AI label'io, tai sunku pasakyti.*“ Ši pastaba rodo, kad dalis vartotojų, susidūrę su universiteto komunikacija, tikisi aiškaus ir atviro žymėjimo apie DI naudojimą.

Apibendrinant galima teigti, kad vartotojų emocinis atsakas į socialinių tinklų įrašus stipriai priklauso nuo turinio kokybės, nuoseklumo, kalbinės raiškos ir pateikimo skaidrumo. Neigiamos reakcijos daugiausia kilo iš tekstinių ir vizualinių neatitikimų, neutralūs atsakymai atspindėjo abejones ir vertinimo sudėtingumą, o teigiami – buvo susiję su profesionaliu ir natūraliu turinio pateikimu.

4.5. Hipotezių tikrinimo rezultatai

Nepriklausomų imčių t-testas. Nepriklausomų imčių t-testas (angl. Independent Samples t-test). buvo atliktas siekiant patikrinti pirmąją tyrimo hipotezę:

H1: DI sukurtas ir pažymėtas rinkodaros turinys (palyginti su nepažymėtu) sukels didesnę suvokiamą turinio ir jo šaltinio atitikimą.

Šis analizės būdas pasirinktas todėl, kad leidžia įvertinti vidurkių skirtumus tarp dviejų nepriklausomų grupių – respondentų, mačiusių turinį su DI autorystės žyma, ir tų, kurie matė nepažymėtą turinį. Ši analizė padeda nustatyti, ar turinio žymėjimas kaip sukurtas DI turi statistiškai reikšmingą poveikį vartotojų suvokimui apie tai, kiek turinys atitinka jį išpublikavusį šaltinį.

Analizės rezultatai (žr. 24 lentelė) parodė, kad respondentai, kurie matė turinį pažymėtą kaip sukurtą DI, šaltinio ir turinio atitikimą įvertino aukščiau ($M = 4,58$, $SD = 1,09$), nei tie, kurie matė nepažymėtą turinį ($M = 4,16$, $SD = 1,12$).

24 lentelė. Suvokiamo šaltinio ir turinio atitikimo vidurkiai pagal DI autorystės žymėjimą (N = 372)

Pažymėta, kad įrašą sukūrė DI	N	Vidurkis (M)	Standartinis nuokrypis (SD)
Ne	210	4,1619	1,12306
Taip	162	4,5802	1,09302

Prieš atlikdamas t-testą, buvo patikrinta dispersijų lygybė tarp grupių naudojant Levene'o testą (žr. 22 priedas). Rezultatai parodė, kad dispersijos statistiškai reikšmingai nesiskiria ($p=0,963$), todėl atlikta analizė, laikantis lygių dispersijų prielaidos. Nepriklausomų imčių t-testas parodė, kad šis skirtumas yra statistiškai reikšmingas: $t(370) = -3,604$, $p < 0,001$. Tai rodo, kad autorystės žyma turėjo

reikšmingą poveikį respondentų suvokimui – pažymėtas turinys labiau atitiko jų lūkesčius apie suderinamumą tarp žinutės turinio ir ją pateiklusio šaltinio.

Norint įvertinti šio skirtumo praktinę reikšmę, buvo apskaičiuoti efektų dydžiai. Pagal Cohen'o (1988) interpretaciją šios reikšmės rodo vidutinio stiprumo efektą. Tai reiškia, kad DI autorystės žymėjimas turėjo statistiškai reikšmingą poveikį vartotojų suvokiamam šaltinio ir turinio atitikimui, tačiau šio poveikio stiprumas nėra didelis – jis vertinamas kaip vidutinis. Šie rezultatai patvirtina, kad DI autorystės žymėjimas daro reikšmingą poveikį vartotojų suvokiamam šaltinio ir turinio atitikimui. Respondentai, kurie matė turinį su aiškiai pažymėta DI autoryste, vertino jį kaip geriau atitinkantį jo šaltinį nei tie, kurie matė nepažymėtą turinį. Nors vidurkių skirtumas nėra didelis, statistinė analizė parodė, kad šis skirtumas yra reikšmingas, o efektas – vidutinio stiprumo. Tai leidžia daryti išvadą, kad autorystės žyma gali sustiprinti vartotojų suvokimą apie informacijos suderinamumą su jos šaltiniu. Tokie rezultatai atitinka schemos atitikimo teorijos prielaidas, teigiančias, jog informacija, aiškiai perteikianti savo kilmę, yra lengviau įsisavinama ir vertinama kaip nuoseklesnė. **Todėl hipotezė H1 – jog pažymėtas DI sukurtas turinys skatina didesnę šaltinio ir turinio atitikimo suvokimą nei nepažymėtas – yra patvirtinama.**

Dvifaktorinė dispersijos analizė. Dvifaktorinė nepriklausomų imčių dispersinė analizė (angl. Two-way ANOVA) buvo atlikta siekiant patikrinti antrąją tyrimo hipotezę:

H2: DI sukurtas ir pažymėtas rinkodaros turinys (palyginti su nepažymėtu) sukels didesnę suvokiamą turinio ir jo šaltinio atitikimą kai turinys turi sąsają su DI.

Analizėje buvo įtraukti du nepriklausomi kintamieji: DI autorystės žymėjimas (pažymėta ar ne) ir turinio sąsaja su DI (yra ar nėra), o priklausomas kintamasis – suvokiamas šaltinio ir turinio atitikimas. Levene'o testas (žr. 25 lentelė) parodė, kad duomenų dispersijos tarp grupių nesiskiria statistiškai reikšmingai ($p=0,074$), todėl tenkinta dispersijų lygybės prielaida, leidžianti pagrįstai taikyti ANOVA.

25 lentelė. Lygių dispersijų prielaidos Levene'o testas: DI autorystės žymėjimo ir turinio sąsajos su DI poveikis suvokiamam šaltinio ir turinio atitikimui ($N = 372$)

Vertinimo pagrindas	Levene'o statistika	Laisvės laipsniai (df1)	Laisvės laipsniai (df2)	p reikšmė (Sig.)
Pagal vidurkį	2,328	3	368	0,074
Pagal medianą	2,007	3	368	0,113
Pagal medianą (pakoreguoti df)	2,007	3	350,182	0,113
Pagal nukirptą vidurkį	2,358	3	368	0,071

Dvifaktorinės dispersijos analizės rezultatai (žr. 26 lentelė) atskleidė statistiškai reikšmingą pagrindinį efektą tiek autorystės žymėjimui, $F(1, 368) = 9,320$, $p=0,002$, tiek turinio sąsajai su DI, $F(1, 368) = 8,822$, $p=0,003$. Svarbiausia, kad taip pat nustatytas statistiškai reikšmingas sąveikos efektas, $F(1, 368) = 4,746$, $p=0,030$. Šis rezultatas leidžia teigti, kad autorystės žymėjimo poveikis priklausė nuo to, ar turinys buvo susijęs su DI.

26 lentelė. Dvifaktorinės dispersijos analizės rezultatai: DI žymėjimo ir turinio sąsajos su DI poveikis suvokiamam šaltinio ir turinio atitikimui ($N = 372$)

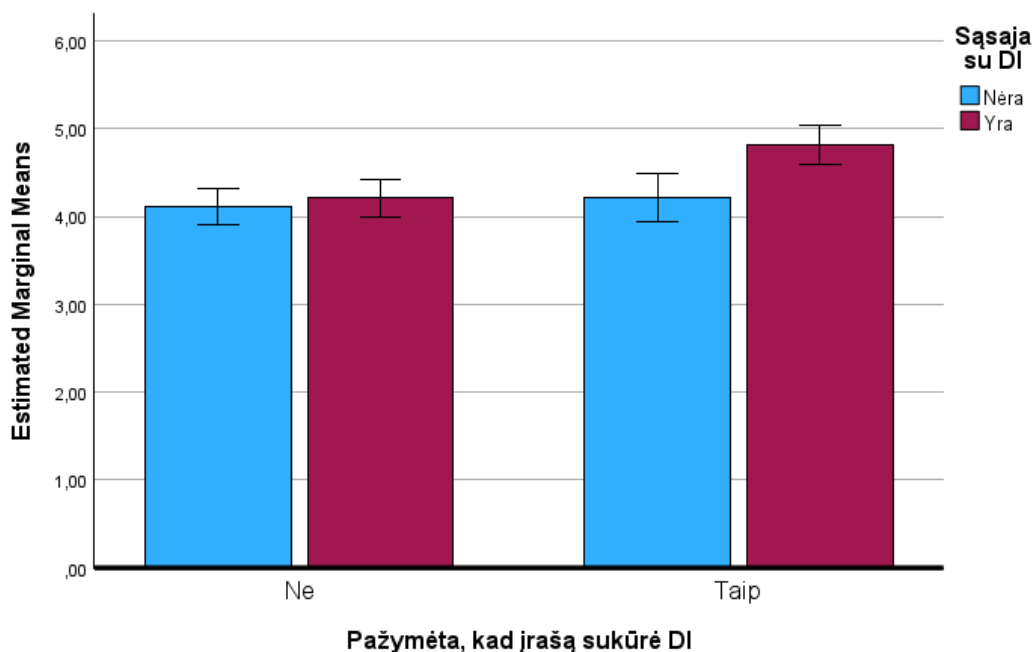
Šaltinis	III tipo kvadratų suma	Laisvės laipsniai (df)	Vidutinis kvadratas	F reikšmė	p reikšmė (Sig.)	Dalinis eta kvadratas
Koreguotas modelis	30,273	3	10,091	8,408	<0,001	0,064
Konstanta	6714,343	1	6714,343	5594,223	<0,001	0,938
Autorystės žyma	11,186	1	11,186	9,320	0,002	0,025
Sąsaja su DI	10,588	1	10,588	8,822	0,003	0,023
Autorystės žymos ir sąsajos su DI sąveika	5,697	1	5,697	4,746	0,030	0,013
Pastaba. $R^2 = 0,64$ (koreguotas $R^2 = 0,057$)						

Rezultatų palyginimai skirtingose scenarijų grupėse (žr. 27 lentelė) parodė, kad didžiausias šaltinio ir turinio atitikimo vertinimo vidurkis buvo tarp tų respondentų, kurie matė pažymėtą įrašą, pristatantį inovaciją, susijusią su DI ($M = 4,8163$, $SD = 0,95$). Tuo tarpu mažiausi vidurkiai fiksuoti tarp tų, kurie matė nepažymėtą įrašą, kurio turinys nebuvo susijęs su DI ($M = 4,1173$, $SD = 1,05$). Tai leidžia manyti, kad autorystės žymėjimo poveikis suvokiamam šaltinio ir turinio atitikimui priklauso nuo to, ar turinys yra susijęs su DI – pažymėtas DI turinys, kai jis susijęs su DI tematika, vertinamas palankiausiai.

27 lentelė. Suvokiamo šaltinio ir turinio atitikimo vidurkiai pagal DI žymėjimą ir turinio sąsają su DI ($N = 372$)

Pažymėta, kad įrašą sukūrė DI	Sąsaja su DI	Vidurkis (M)	Standartinis nuokrypis (SD)	N
Ne	Nėra	4,1173	1,05342	108
	Yra	4,2092	1,19581	102
	Iš viso	4,1619	1,12306	210
Taip	Nėra	4,2188	1,20730	64
	Yra	4,8163	0,94538	98
	Iš viso	4,5802	1,09302	162
Iš viso	Nėra	4,1550	1,11076	172
	Yra	4,5067	1,11982	200
	Iš viso	4,3441	1,12788	372

Toliau pateikiami vizualiai pavaizduoti vidutiniai šaltinio ir turinio atitikimo įvertinimai pagal dvi nepriklausomų imčių dispersijos analizėje tirtas sąlygas (5 pav.)



5 pav. Šaltinio ir turinio atitikimo vertinimų palyginimai nepriklausomų kintamųjų grupėse (N = 372)

Šie rezultatai leidžia daryti išvadą, kad tiek autorystės žymėjimas, tiek turinio sąsaja su DI turi nepriklausomą statistiškai reikšmingą poveikį suvokiamam šaltinio ir turinio atitikimui. Svarbiausia, jog tarp šių dviejų kintamųjų taip pat buvo nustatytas statistiškai reikšmingas sąveikos efektas ($F(1, 368) = 4,746, p=0,030$), rodantis, kad autorystės žymėjimo poveikis priklausė nuo to, ar turinys buvo susijęs su DI. Kitaip tariant, žymėjimo įtaka nebuvo vienoda visais atvejais – ji sustiprėjo būtent tada, kai pats turinys turėjo aiškią teminę sąsają su dirbtiniu intelektu. Šį poveikį papildomai iliustruoja vidurkiai skirtingose eksperimentinėse grupėse (žr. 27 lentelė): didžiausias šaltinio ir turinio atitikimas fiksuotas tuomet, kai turinys buvo pažymėtas kaip sukurtas DI ir turėjo teminę sąsają su DI, o žemiausias – kai nebuvo nei pažymėjimo, nei teminio ryšio. Tai patvirtina, kad žymėjimo poveikis nėra vienodas visose situacijose – jis sustiprėja tik specifiniuose kontekstuose, priklausomai nuo turinio pobūdžio.

Dalinių eta kvadratų (η^2) reikšmės rodo, kad žymėjimo efektas paaiškino 2,5 proc. dispersijos, turinio sąsajos efektas – 2,3 proc., o sąveikos efektas – 1,3 proc. Nors šie efektai nėra labai dideli, jie yra statistiškai reikšmingi ir patvirtina, kad abiejų kintamųjų sąveika daro įtaką respondentų vertinimams. **Todėl, remiantis šia analize, H2 hipotezė yra patvirtinama.**

Regresinė analizė. Regresinė analizė buvo atlikta siekiant patikrinti kitas tyrimo hipotezes:

H3a: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų suvokiamą šaltinio kompetenciją.

H3b: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų suvokiamą šaltinio empatiją.

H3c: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų suvokiamą šaltinio palankumą.

H3d: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų pasitikėjimą šaltiniu.

Atliekant paprastąją tiesinę regresiją, priklausomas kintamasis buvo **šaltinio kompetencijos įvertinimas**, o nepriklausomas – suvokiamas šaltinio ir turinio atitikimas. Modelio tinkamumas buvo įvertintas per F-testą ir determinacijos koeficientus. Rezultatai rodo (žr. 28 lentelė), jog modelis yra statistiškai reikšmingas: $F(1, 370) = 91,575$, $p < 0,001$, o determinacijos koeficiento R^2 reikšmė – 0,198, o tai reiškia, kad šaltinio ir turinio atitikimas paaiškina 19,8 proc. šaltinio kompetencijos įvertinimo dispersijos. Koreguota R^2 reikšmė siekia 0,196.

28 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis šaltinio kompetencijai (N = 372)

R	R^2	Koreguotas R^2	Standartinė paklaida	R^2 pokytis	F pokytis	Laisvės laipsniai (df1)	Laisvės laipsniai (df2)	p reikšmė
0,445	0,198	0,196	1,11472	0,198	91,575	1	370	<0,001

Regresijos koeficientų analizės rezultatai (žr. 29 lentelė) atskleidė, kad nepriklausomas kintamasis – suvokiamas šaltinio ir turinio atitikimas – yra statistiškai reikšmingas prognozuojant šaltinio kompetenciją. Standartizuotas regresijos koeficientas (β) siekė 0,445, o t reikšmė – 9,569, $p < 0,001$. Šios reikšmės rodo, kad tarp šaltinio ir turinio atitikimo suvokimo ir šaltinio kompetencijos vertinimo egzistuoja teigiamas ir vidutinio stiprumo statistiškai reikšmingas ryšys. Teigiamas koeficientas reiškia, jog augant respondentų įsitikinimui, kad turinys atitinka jį pateikusį šaltinį, didėja ir šaltinio kompetencijos įvertinimas. Tai leidžia daryti išvadą, kad vartotojai, kurie socialinių tinklų įrašė mato aiškia vidinę darną tarp žinutės turinio ir komunikacijos šaltinio, yra labiau linkę manyti, jog šaltinis yra kompetentingas.

29 lentelė. Regresijos koeficientų analizės rezultatai, tiriant šaltinio ir turinio atitikimo poveikį šaltinio kompetencijai (N = 372)

Kintamasis	B reikšmė	Standartinė paklaida	Standartizuotas Beta	t reikšmė	p reikšmė
Šaltinio ir turinio atitikimas	0,491	0,051	0,445	9,569	<0,001

Galiausiai, galima daryti išvadą, kad suvoktas atitikimas tarp šaltinio ir turinio reikšmingai prisideda prie aukštesnio šaltinio kompetencijos vertinimo formavimo.

Toliau buvo atlikta paprastoji tiesinė regresija, kurios pagalba buvo įvertintas ryšys tarp suvokiamo šaltinio ir turinio atitikimo bei vienos iš šaltinio patikimumo dimensijų – **šaltinio palankumo**. Modelio suvestinės analizė (žr. 30 lentelė) rodo, kad modelis yra statistiškai reikšmingas: $F(1, 370) = 40,662$, $p < 0,001$. Determinacijos koeficientas (R^2) siekia 0,099, o pakoreguotas R^2 – 0,097, t. y. modelis paaiškina apie 9,9 % šaltinio palankumo įvertinimo dispersijos. Nors prognozinė modelio kokybė nėra aukšta, jis laikomas prasmingu komunikacijos ir elgsenos tyrimų kontekste, kuriame dažnai pasitaiko daugybė įtakos veiksnių.

30 lentelė. Modelio suvestinė: šaltinio ir turinio atitikimo poveikis šaltinio palankumui (N = 372)

R	R^2	Koreguotas R^2	Standartinė paklaida	R^2 pokytis	F pokytis	Laisvės laipsniai (df1)	Laisvės laipsniai (df2)	p reikšmė
0,315	0,099	0,097	1,34659	0,099	40,662	1	370	<0,001

Regresijos koeficientų analizė (žr. 31 lentelė) patvirtino, kad šaltinio ir turinio atitikimas yra reikšmingas palankumo prognozuotojas ($\beta = 0,315$, $t = 6,377$, $p < 0,001$). Šis teigiamas poveikis rodo, kad kuo labiau vartotojai suvokia turinio ir jo šaltinio atitikimą, tuo labiau jie linkę vertinti šaltinį kaip patikimą.

31 lentelė. Regresijos koeficientų analizės rezultatai: šaltinio ir turinio atitikimo poveikis šaltinio palankumui (N = 372)

Kintamasis	B reikšmė	Standartinė paklaida	Standartizuotas Beta	t reikšmė	p reikšmė
Šaltinio ir turinio atitikimas	0,395	0,062	0,315	6,377	<0,001

Šie rezultatai patvirtina, kad suvoktas šaltinio ir turinio atitikimas turi statistiškai reikšmingą teigiamą poveikį šaltinio palankumo vertinimams.

Toliau atlikta paprastosios tiesinės regresijos analizė, kuria buvo siekiama įvertinti ryšį tarp suvokiamo šaltinio ir turinio atitikimo bei vienos iš šaltinio patikimumo dimensijų – **šaltinio empatijos**. Modelio suvestinės rezultatai (žr. 32 lentelė) rodo, kad regresinis modelis yra statistiškai reikšmingas: $F(1, 370) = 99,761$, $p < 0,001$. Determinacijos koeficientas $R^2 = 0,212$ leidžia teigti, kad modelis paaiškina apie 21,2 proc. šaltinio empatijos vertinimo dispersijos, o pakoreguotas $R^2 = 0,210$ rodo šio įvertinimo stabilumą.

32 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis šaltinio empatijai (N = 372)

R	R ²	Koreguotas R ²	Standartinė paklaida	R ² pokytis	F pokytis	Laisvės laipsniai (df1)	Laisvės laipsniai (df2)	p reikšmė
0,461	0,212	0,210	1,12648	0,212	99,761	1	370	<0,001

Regresijos koeficientų analizė (žr. 33 lentelė) atskleidė, kad šaltinio ir turinio atitikimas yra reikšmingas prognozuojant šaltinio empatiją ($\beta = 0,461$, $t = 9,988$, $p < 0,001$). Šis rezultatas rodo teigiamą ryšį tarp kintamųjų: kuo vartotojai labiau įžvelgia šaltinio ir turinio atitikimą, tuo stipresnis yra jų šaltinio empatijos įvertinimas. Standartizuota koeficiento reikšmė ($\beta = 0,461$), o tai leidžia daryti išvadą, kad atitikimo suvokimas atlieka svarbų vaidmenį formuojant vartotojų empatiją šaltinio atžvilgiu.

33 lentelė. Regresijos koeficientų analizės rezultatai, tiriant šaltinio ir turinio atitikimo poveikį šaltinio empatijai (N = 372)

Kintamasis	B reikšmė	Standartinė paklaida	Standartizuotas Beta	t reikšmė	p reikšmė
Šaltinio ir turinio atitikimas	0,518	0,052	0,461	9,988	<0,001

Šaltinio pasitikėjimo dimensijai buvo taikyta paprastoji tiesinė regresijos analizė, siekiant įvertinti, ar suvokiamas šaltinio ir turinio atitikimas lemia aukštesnius šaltinio pasitikėjimo vertinimus.

Rezultatai (žr. 34 lentelė) atskleidė, kad modelis yra statistiškai reikšmingas: $F(1,370) = 59,735$, $p < 0,001$. Determinacijos koeficientas (R^2) siekia 0,139, o pakoreguotas $R^2 = 0,137$, tai reiškia, kad šaltinio ir turinio atitikimo suvokimas paaiškina apie 13,9 proc. šaltinio pasitikėjimo vertinimo dispersijos.

34 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis šaltinio pasitikėjimui (N = 372)

R	R ²	Koreguotas R ²	Standartinė paklaida	R ² pokytis	F pokytis	Laisvės laipsniai (df1)	Laisvės laipsniai (df2)	p reikšmė
0,373	0,139	0,137	1,26795	0,139	59,735	1	370	<0,001

Regresijos koeficientų analizė (žr. 35 lentelė) rodo, kad šaltinio ir turinio atitikimas yra statistiškai reikšmingas šaltinio pasitikėjimo prediktorius ($\beta = 0,373$, $t = 7,729$, $p < 0,001$). Teigiama regresijos koeficiento kryptis rodo, kad didėjant atitikimo tarp šaltinio ir turinio suvokimui, didėja ir respondentų šaltinio pasitikėjimo vertinimai. Tai reiškia, jog nuoseklumas tarp pateikiamo turinio ir jo šaltinio padeda stiprinti pasitikėjimą šaltiniu.

35 lentelė. Regresijos koeficientų analizė: šaltinio ir turinio atitikimo poveikis šaltinio pasitikėjimui (N = 372)

Kintamasis	B reikšmė	Standartinė paklaida	Standartizuotas Beta	t reikšmė	p reikšmė
Šaltinio ir turinio atitikimas	0,451	0,058	0,373	7,729	<0,001

Šie rezultatai leidžia teigti, kad šaltinio ir turinio atitikimas yra svarbus veiksnys, formuojantis vartotojų pasitikėjimą informacijos šaltiniu. Kai turinys yra suvokiamas kaip tinkantis ir nuosekliai susijęs su jį pateikiančiu šaltiniu, didėja tikimybė, kad šaltinis bus laikomas patikimu.

Atliekant atskiras regresines analizes kiekvienai šaltinio patikimumo dimensijai – kompetencijai, palankumui, empatijai ir pasitikėjimui – visais atvejais buvo nustatytas statistiškai reikšmingas teigiamas ryšys tarp suvokiamo šaltinio ir turinio atitikimo. Kuo vartotojai geriau įžvelgė šaltinio ir turinio dermę, tuo aukštesni buvo jų šaltinio patikimumo įvertinimai visose tirtose dimensijose. Šie rezultatai parodo, kad vartotojų suvokiamas šaltinio ir turinio atitikimas yra nuoseklus ir reikšmingas prediktorius visoms šaltinio patikimumo dimensijoms. **Tai leidžia pagrįstai teigti, kad hipotezės H3a, H3b, H3c ir H3d – teigiančios apie ryšius tarp suvokto atitikimo ir šaltinio patikimumo – yra empiriškai pagrįstos ir statistiškai patvirtintos.**

Paprastoji tiesinė regresinė analizė buvo taikyta tikrinant ketvirtąjį tyrimo hipotezę:

H4: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų suvokiamą turinio patikimumą.

Šiai hipotezei patikrinti turinio patikimumas buvo traktuojamas kaip priklausomas kintamasis, o šaltinio ir turinio atitikimas – kaip nepriklausomas prognozuojantis kintamasis. Toks analizės metodas buvo pasirinktas siekiant įvertinti, ar aukštesnis šaltinio ir turinio atitikimo suvokimas lemia didesnę vartotojų pasitikėjimą pačiu turiniu.

Regresinės analizės modelio santrauka (žr. 36 lentelė) rodo, kad regresinis modelis yra statistiškai reikšmingas – $F(1, 370) = 93,783$, $p < 0,001$. Determinacijos koeficientas (R^2) siekia 0,202, o koreguotas R^2 – 0,200, o tai reiškia, kad šaltinio ir turinio atitikimo suvokimas paaiškina 20,2 proc. turinio patikimumo vertinimo dispersijos. Tai yra gana reikšminga proporcija, ypač atsižvelgiant į socialinių mokslų tyrimų kontekstą.

36 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis turinio patikimumui (N = 372)

R	R ²	Koreguotas R ²	Standartinė paklaida	R ² pokytis	F pokytis	Laisvės laipsniai (df1)	Laisvės laipsniai (df2)	p reikšmė
0,450	0,202	0,200	1,25134	0,202	93,783	1	370	<0,001

Regresijos koeficientų analizės duomenys (žr. 37 lentelė) rodo, kad šaltinio ir turinio atitikimas yra reikšmingas turinio patikimumo prediktorius: $\beta = 0,450$, $t = 9,684$, $p < 0,001$. Numatomas teigiamas poveikis rodo, kad vartotojų suvokiamas šaltinio ir turinio atitikimas yra didesnis, tuo patikimesnis jiems atrodo pats turinys.

37 lentelė. Regresijos koeficientų analizės rezultatai: šaltinio ir turinio atitikimo poveikis turinio patikimumui (N = 372)

Kintamasis	B reikšmė	Standartinė paklaida	Standartizuotas Beta	t reikšmė	p reikšmė
Šaltinio ir turinio atitikimas	0,558	0,058	0,450	9,684	<0,001

Remiantis šia analize, galima daryti išvadą, kad šaltinio ir turinio atitikimo suvokimas reikšmingai prisideda prie vartotojų vertinamo turinio patikimumo formavimo. Taigi, **hipotezė H4 yra pagrįsta empiriniais duomenimis ir patvirtinama.**

Paprastoji tiesinė regresijos analizė buvo atlikta siekiant patikrinti penktąją tyrimo hipotezę:

H5: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų įsitraukimo ketinimus socialiniuose tinkluose.

Ja siekta įvertinti, ar stipresnis vartotojų suvokimas, kad turinys dera su jį pateikusių šaltiniu, prognozuoja didesnę ketinimą įsitraukti į tokį turinį (pavyzdžiui, jį komentuoti, juo dalintis ar spausti „patinka“). Įsitraukimo ketinimai tyrime buvo vertinti kaip priklausomas kintamasis, o suvokiamas šaltinio ir turinio atitikimas – kaip nepriklausomas prognozuojantis veiksnys.

Modelio tinkamumas buvo įvertintas per F-testą ir determinacijos koeficiento (R^2) rodiklius. Kaip matyti iš 38 lentelės, regresijos modelis yra statistiškai reikšmingas: $F(1, 370) = 18,173$, $p < 0,001$, tačiau tai, jog determinacijos koeficientas (R^2) lygus 0,047, o pakoreguotas R^2 siekia 0,044, rodo, kad šaltinio ir turinio atitikimo suvokimas paaiškina tik 4,7 proc. įsitraukimo ketinimų dispersijos. Todėl modelis laikytinas nepakankamos kokybės.

38 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis vartotojų įsitraukimo ketinimams (N = 372)

R	R ²	Koreguotas R ²	Standartinė paklaida	R ² pokytis	F pokytis	Laisvės laipsniai (df1)	Laisvės laipsniai (df2)	p reikšmė
0,216	0,047	0,044	1,32119	0,047	18,173	1	370	<0,001

Toliau pateikta regresijos koeficientų lentelė (žr. 39 lentelė) rodo, kad šaltinio ir turinio atitikimas yra reikšmingas prediktorius įsitraukimo ketinimams ($\beta = 0,216$, $t = 4,263$, $p < 0,001$). Teigiamas regresijos koeficientas ($B = 0,259$) leidžia teigti, kad didesnis atitikimo suvokimas yra susijęs su didesniu noru sąveikauti su turiniu socialiniuose tinkluose.

39 lentelė. Regresijos koeficientų analizė: šaltinio ir turinio atitikimo poveikis įsitraukimo ketinimams (N = 372)

Kintamasis	B reikšmė	Standartinė paklaida	Standartizuotas Beta	t reikšmė	p reikšmė
Šaltinio ir turinio atitikimas	0,259	0,061	0,216	4,263	<0,001

Atsižvelgiant į šiuos rezultatus, H5 hipotezė nėra patvirtinama. Nors rezultatai rodo statistiškai reikšmingą ryšį tarp vartotojų suvokiamo šaltinio ir turinio atitikimo bei jų įsitraukimo ketinimų socialiniuose tinkluose, poveikio dydis yra labai mažas ($R^2 = 0,047$), paaiškinantis tik nedidelę dalį įsitraukimo ketinimų dispersijos. Tai reiškia, kad nors ryšys egzistuoja, jis yra silpnas ir praktiniu požiūriu riboto reikšmingumo.

Norint įvertinti, ar vartotojų suvokiamas šaltinio ir turinio patikimumas prognozuoja jų ketinimus įsitraukti į turinį socialiniuose tinkluose, buvo atlikta daugialypė tiesinė regresinė analizė. Šiuo atveju tikrintos penkios hipotezės:

- H6a: Vartotojų suvokiama šaltinio kompetencija teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;
- H6b: Vartotojų suvokiama šaltinio empatija teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;
- H6c: Vartotojų suvokiamas šaltinio palankumas teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;
- H6d: Vartotojų suvokiamas pasitikėjimas šaltiniu teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;
- H7: Vartotojų suvokiamas turinio patikimumas teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose.

Analizės metu buvo įtraukti penki nepriklausomi kintamieji: turinio patikimumas, šaltinio kompetencija, palankumas, empatija ir pasitikėjimas. Priklausomas kintamasis – vartotojų įsitraukimo ketinimai.

Modelio suvestinės rezultatai (žr. 40 lentelė) rodo, kad regresinis modelis yra statistiškai reikšmingas: $F(5, 366) = 29,448$, $p < 0,001$. Determinacijos koeficientas (R^2) siekia 0,287, o pakoreguotas R^2 – 0,277, o tai reiškia, kad modelis paaiškina apie 28,7 proc. vartotojų įsitraukimo ketinimų dispersijos. Tokia vertė laikytina pakankamai reikšminga, atsižvelgiant į socialinių mokslų kontekstą, kuriame įsitraukimo elgsena paprastai priklauso nuo daugybės tarpusavyje susijusių veiksnių.

40 lentelė. Modelio santrauka: šaltinio ir turinio atitikimo poveikis vartotojų įsitraukimo ketinimams (N = 372)

R	R^2	Koreguotas R^2	Standartinė paklaida	R^2 pokytis	F pokytis	Laisvės laipsniai (df1)	Laisvės laipsniai (df2)	p reikšmė
0,536	0,287	0,277	1,14899	0,287	29,448	5	366	<0,001

Regresijos koeficientų analizė (žr. 41 lentelė) parodė, kad statistiškai reikšmingą teigiamą poveikį įsitraukimo ketinimams turėjo turinio patikimumas ($\beta = 0,303$, $p < 0,001$), šaltinio palankumas ($\beta = 0,253$, $p < 0,001$) ir šaltinio empatija ($\beta = 0,337$, $p < 0,001$). Tai reiškia, kad kuo vartotojams turinys atrodo patikimesnis, o jo šaltinis – empatiškas ir palankus, tuo labiau jie linkę įsitraukti į tokį turinį socialinių tinklų platformose. Šis ryšys gali pasireikšti per aktyvų elgseninį įsitraukimą – reakcijas, komentarus ar dalijimąsi turiniu.

Tuo tarpu pasitikėjimas šaltiniu turėjo neigiamą poveikį įsitraukimui ($\beta = -0,278$, $p < 0,001$), o šaltinio kompetencija – statistiškai nereikšmingą ($p = 0,235$). Šis rezultatas gali rodyti, kad didelis pasitikėjimas šaltiniu nebūtinai skatina įsitraukimą, o veikiau gali būti susijęs su pasyvesniu informacijos priėmimu. Tuo tarpu kompetencijos suvokimas gali būti svarbesnis kitiems vertinimo aspektams, bet nebūtinai motyvuoja elgseninius veiksmus.

41 lentelė. Pagrindiniai rezultatai: koeficientų analizė ($N = 372$)

Kintamasis	B reikšmė	Standartinė paklaida	Standartizuotas Beta	t reikšmė	p reikšmė	Tolerancija	Dispersijos didinimo faktoriai (VIF)
Turinio patikimumas	0,292	0,059	0,303	4,940	<0,001	0,519	1,927
Šaltinio kompetencija	-0,086	0,072	-0,079	-1,189	0,235	0,445	2,249
Šaltinio palankumas	0,241	0,051	0,253	4,747	<0,001	0,686	1,459
Šaltinio empatija	0,359	0,066	0,337	5,402	<0,001	0,502	1,992
Pasitikėjimas šaltiniu	-0,275	0,064	-0,278	-4,323	<0,001	0,471	2,123

Multikolinearumo patikrinimui buvo įvertinti tolerancijos ir dispersijos didinimo faktoriai (VIF). Visi VIF rodikliai buvo mažesni nei 2,5, o tolerancijos reikšmės – virš 0,4, todėl galima pagrįstai teigti, kad modelyje nėra reikšmingo multikolinearumo. Tai reiškia, kad kiekvienas prediktorius prisideda prie modelio paaiškinamosios galios nepriklausomai nuo kitų kintamųjų.

Apibendrinant galima teigti, kad **šeštosios b ir c hipotezės (H6b ir H6c) patvirtinamos – šaltinio empatija ir palankumas turėjo teigiamą poveikį, tačiau kitų dimensijų hipotezės (H6a ir H6d) – ne, mat kompetencija neturėjo reikšmingo poveikio, o pasitikėjimas turėjo neigiamą**. Septintoji **hipotezė (H7) yra patvirtinama** – turinio patikimumas turi statistiškai reikšmingą teigiamą poveikį įsitraukimo ketinimams. Šie rezultatai rodo, kad šaltinio ir turinio patikimumo poveikis nėra vienalytis, todėl svarbu turinio patikimumą ir šaltinio patikimumo dimensijas atskirai.

Siekiant įvertinti, kaip vartotojų suvokiamas šaltinio ir turinio atitikimas prognozuoja jų ketinimus įsitraukti į turinį socialiniuose tinkluose per turinio ir šaltinių patikimumo atskirus aspektus, buvo atlikta mediavimo analizė naudojant specializuotą makrokomandą *PROCESS* (Hayes, 2013). Buvo tikrinamos šios hipotezės:

- *H8a: Šaltinio kompetencija teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose;*
- *H8b: Šaltinio empatija teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose;*
- *H8c: Šaltinio palankumas teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose;*
- *H8d: Pasitikėjimas šaltiniu teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose.*
- *H9: Turinio patikimumas teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose.*

Regresinių analizių rezultatai pateikiami lentelėse: 42, 43, 44, 45, 46, 47, 48, 49, 50 ir 51.

42 lentelė. Regresijos modelių, tiriant ryšį tarp šaltinio ir turinio atitikimo, turinio patikimumo ir įsitraukimo ketinimų, rezultatai

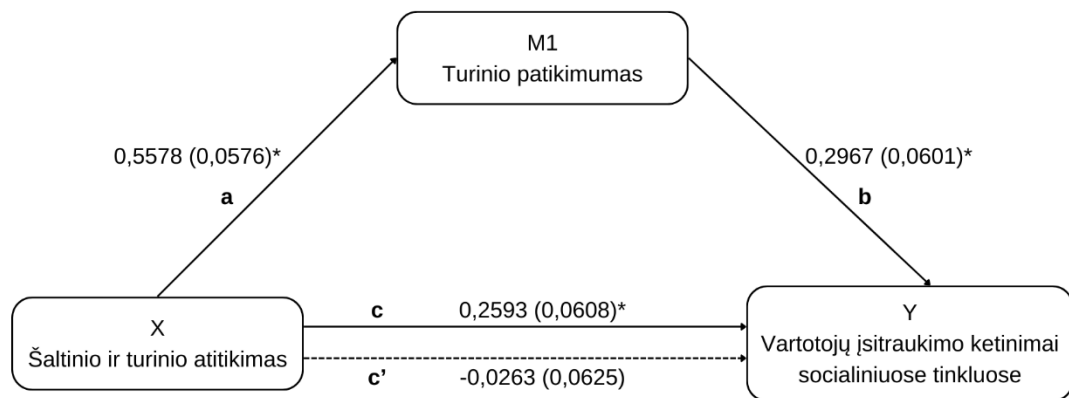
R	Priklausomas kintamasis											
	M1: Turinio patikimumas				Y: Įsitraukimo ketinimai				Y: Įsitraukimo ketinimai			
	→	Koef.	SE	p	→	Koef.	SE	p	→	Koef.	SE	p
Konstanta	iM →	1,7614	0,2585	0,0000	iY →	0,5751	0,2866	0,0455	iY →	1,5064	0,2729	0,0000
X: Šaltinio ir turinio atitikimas	a →	0,5578	0,0576	0,0000	c' →	-0,0263	0,0625	0,6743	c →	0,2593	0,0608	0,0000
M1: Turinio patikimumas	-	-	-	-	b →	0,2967	0,0601	0,0000	-	-	-	-
Modelis	$R^2 = 0,2022$, $F(1;370) = 93,7827$, $p=0,0000$				$R^2 = 0,2872$, $F(6;365) = 24,5143$, $p=0,0000$				$R^2 = 0,0468$, $F(1;370) = 18,1735$, $p=0,0000$			

Remiantis pateiktais duomenimis (žr. 42 lentelė) galima teigti, kad egzistuoja statistiškai reikšminga šaltinio ir turinio atitikimo įtaka turinio patikimumui (a kelias; $p=0,000$) ir įsitraukimo ketinimams (b kelias; $p=0,000$). Šaltinio ir turinio atitikimo suminio efekto įtaka įsitraukimo ketinimams irgi statistiškai reikšminga (c kelias; $p=0,000$). 43 lentelėje pateikiama tiesioginė, netiesioginė ir suminė šaltinio ir turinio atitikimo įtaka įsitraukimo ketinimams.

43 lentelė. Tiesioginė, netiesioginė (per turinio patikimumą) ir suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams (N=372)

Kelias	Efektas (EF)	95% pasikliautinis intervalas	
		LLCI	ULCI
Tiesioginė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → įsitraukimo ketinimai (c’)	-0,0263	-0,1492	0,0966
Netiesioginė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → turinio patikimumas → įsitraukimo ketinimai (a*b)	0,1655	0,0920	0,2565
Suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → įsitraukimo ketinimai (c)	0,2593	0,1397	0,3788

Iš rezultatų, pateiktų 43 lentelėje, matoma, kad tarp netiesioginės įtakos (a*b) pasikliautinių intervalų nėra nulio, o tai įrodo, kad turinio patikimumas yra mediatorius ryšyje tarp šaltinio ir turinio atitikimo ir įsitraukimo ketinimų. Atliktos analizės rezultatai pateikiami 6 paveiksle.



6 pav. Šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo kėtinimams per turinio patikimumą

* $p < 0,001$

Modelis patvirtina, kad šaltinio ir turinio atitikimas daro įtaką vartotojų įsitraukimo kėtinimams per turinio patikimumo vertinimą. **Atsižvelgiant į šiuos rezultatus, galima pagrįstai teigti, kad hipotezė H9 yra patvirtinama.**

Remiantis pateiktais regresinės analizės rezultatais, (žr. 44 lentelė) galima teigti, kad šaltinio ir turinio atitikimas turėjo statistiškai reikšmingą poveikį šaltinio kompetencijos vertinimui (a kelias; $p < 0,001$), tačiau šaltinio kompetencija nepasirodė kaip reikšmingas prediktorius vartotojų įsitraukimo kėtinimams (b kelias; $p = 0,2633$). Nepaisant to, bendras šaltinio ir turinio atitikimo poveikis įsitraukimo kėtinimams (c kelias) išliko statistiškai reikšmingas ($p < 0,001$), kas reiškia, kad egzistuoja tiesioginis ryšys tarp šių kintamųjų.

44 lentelė. Regresijos modelių, tiriant ryšį tarp šaltinio ir turinio atitikimo, šaltinio kompetencijos ir įsitraukimo kėtinimų, rezultatai

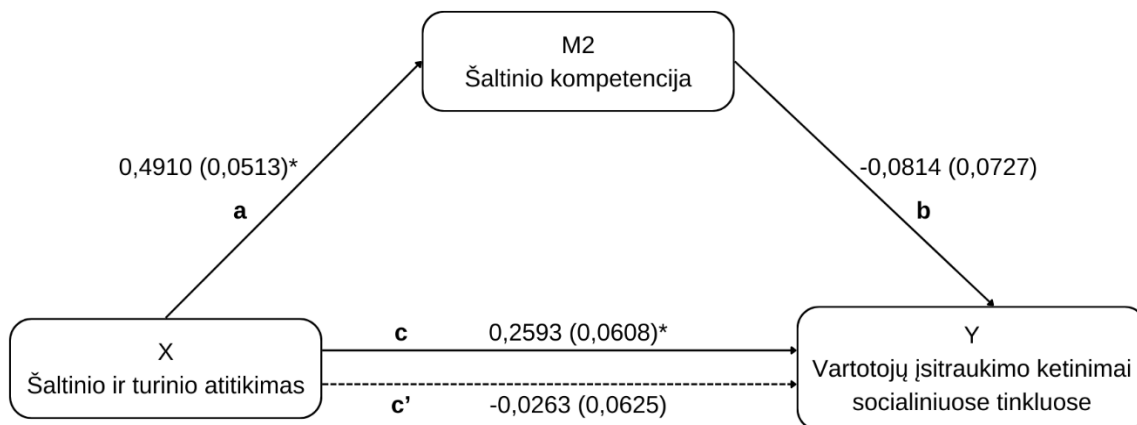
R	Priklausomas kintamasis											
	M2: Šaltinio kompetencija				Y: Įsitraukimo kėtinimai				Y: Įsitraukimo kėtinimai			
	→	Koef.	SE	p	→	Koef.	SE	p	→	Koef.	SE	p
Konstanta	iM →	2,6259	0,2303	0,0000	iY →	0,5751	0,2866	0,0455	iY →	1,5064	0,2729	0,0000
X: Šaltinio ir turinio atitikimas	a →	0,4910	0,0513	0,0000	c'→	-0,0263	0,0625	0,6743	c →	0,2593	0,0608	0,0000
M2: Šaltinio kompetencija	-	-	-	-	b→	-0,0814	0,0727	0,2633	-	-	-	-
Modelis	$R^2 = 0,1984$, $F(1;370) = 91,5745$, $p = 0,0000$				$R^2 = 0,2872$ $F(6;365) = 24,5143$, $p = 0,0000$				$R^2 = 0,0468$ $F(1;370) = 18,1735$, $p = 0,0000$			

45 lentelėje pateikiama tiesioginė, netiesioginė ir suminė šaltinio ir turinio atitikimo įtaka įsitraukimo kėtinimams.

45 lentelė. Tiesioginė, netiesioginė (per šaltinio kompetenciją) ir suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo kėtinimams (N=372)

Kelias	Efektas (EF)	95% pasikliautinis intervalas	
		LLCI	ULCI
Tiesioginė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → įsitraukimo ketinimai (c’)	-0,0263	-0,1492	0,0966
Netiesioginė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → šaltinio kompetencija → įsitraukimo ketinimai (a*b)	-0,0400	-0,1233	0,0270
Suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → įsitraukimo ketinimai (c)	0,2593	0,1397	0,3788

45 lentelėje pateikta netiesioginio poveikio analizė parodė, kad šaltinio kompetencijos įtaka kaip mediatorius nėra statistiškai pagrįsta – pasikliautinis intervalas nuo -0,1233 iki 0,0270 apima nulį, o tai rodo, kad ši sąsaja nėra reikšminga. Vadinas, šaltinio kompetencija nemedijuoja ryšio tarp šaltinio ir turinio atitikimo bei vartotojų įsitraukimo ketinimų. Šie duomenys vizualiai pavaizduoti 7 paveiksle.



7 pav. Šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams per šaltinio kompetenciją

* $p < 0,001$

Atsižvelgiant į tai, hipotezė **H8a**, kurioje teigiama, kad „šaltinio kompetencija teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose“, **nėra patvirtinama**.

Remiantis pateiktais duomenimis (žr. 46 lentelė) galima teigti, kad šaltinio ir turinio atitikimas turėjo statistiškai reikšmingą poveikį šaltinio palankumo vertinimui (a kelias; $p < 0,001$), taip pat ir šaltinio palankumas buvo statistiškai reikšmingas prediktorius vartotojų įsitraukimo ketinimams (b kelias; $p < 0,001$). Nors tiesioginis šaltinio ir turinio atitikimo poveikis įsitraukimo ketinimams (c' kelias; $p = 0,6743$) nebuvo statistiškai reikšmingas, bendras (suminis) poveikis (c kelias; $p < 0,001$) buvo reikšmingas.

46 lentelė. Regresijos modelių, tiriant ryšį tarp šaltinio ir turinio atitikimo, šaltinio palankumo ir įsitraukimo ketinimų, rezultatai

R	Priklausomas kintamasis
---	-------------------------

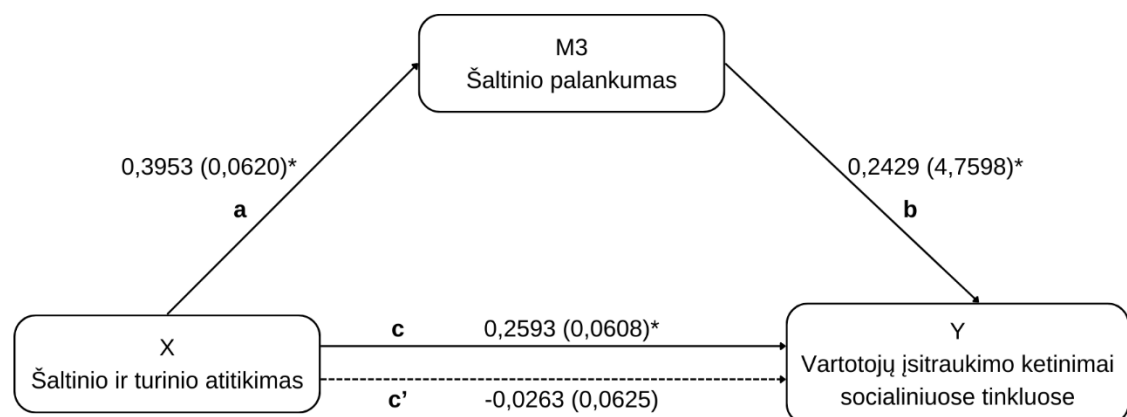
	M3: Šaltinio palankumas				Y: Įsitraukimo ketinimai				Y: Įsitraukimo ketinimai			
	→	Koef.	SE	p	→	Koef.	SE	p	→	Koef.	SE	p
Konstanta	iM →	2,2579	0,2782	0,0000	iY →	0,5751	0,2866	0,0455	iY →	1,5064	0,2729	0,0000
X: Šaltinio ir turinio atitikimas	a →	0,3953	0,0620	0,0000	c'→	-0,0263	0,0625	0,6743	c →	0,2593	0,0608	0,0000
M3: Šaltinio palankumas	-	-	-	-	b→	0,2429	4,7598	0,0000	-	-	-	-
Modelis	$R^2 = 0,0990$, $F(1;370) = 40,6625$, $p=0,0000$				$R^2 = 0,2872$, $F(6;365) = 24,5143$, $p=0,0000$				$R^2 = 0,0468$, $F(1;370) = 18,1735$, $p=0,0000$			

47 lentelėje pateikti netiesioginio poveikio duomenys rodo, kad šaltinio palankumas medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo ketinimų: pasikliautinis intervalas (LLCI = 0,0495; ULCI = 0,1541) neapima nulio, todėl šis ryšys laikomas statistiškai reikšmingu.

47 lentelė. Tiesioginė, netiesioginė (per šaltinio palankumą) ir suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams (N=372)

Kelias	Efektas (EF)	95% pasikliautinis intervalas	
		LLCI	ULCI
Tiesioginė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → įsitraukimo ketinimai (c')	-0,0263	-0,1492	0,0966
Netiesioginė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → šaltinio palankumas → įsitraukimo ketinimai (a*b)	0,0960	0,0495	0,1541
Suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → įsitraukimo ketinimai (c)	0,2593	0,1397	0,3788

Analizės rezultatai vizualizuojami 8 paveiksle.



8 pav. Šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams per šaltinio palankumą

*p<0,001

Atsižvelgiant į rezultatus, galima teigti, kad hipotezė **H8b**, kurioje teigiama, kad „šaltinio palankumas teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose“, **yra patvirtinama**.

Atlikus regresinę analizę (žr. 48 lentelė) paaiškėjo, kad šaltinio ir turinio atitikimas statistiškai reikšmingai prognozavo šaltinio empatijos įvertinimą (a kelias; $p < 0,001$), o šaltinio empatija buvo reikšmingai susijusi ir su įsitraukimo ketinimais (b kelias; $p < 0,001$). Nors tiesioginis šaltinio ir turinio atitikimo poveikis įsitraukimo ketinimams (c' kelias; $p = 0,6743$) nebuvo reikšmingas, suminė šio konstrukto įtaka (c kelias; $p < 0,001$) išliko reikšminga, rodančia bendrą efektą.

48 lentelė. Regresijos modelių, tiriant ryšį tarp šaltinio ir turinio atitikimo, šaltinio empatijos ir įsitraukimo ketinimų, rezultatai

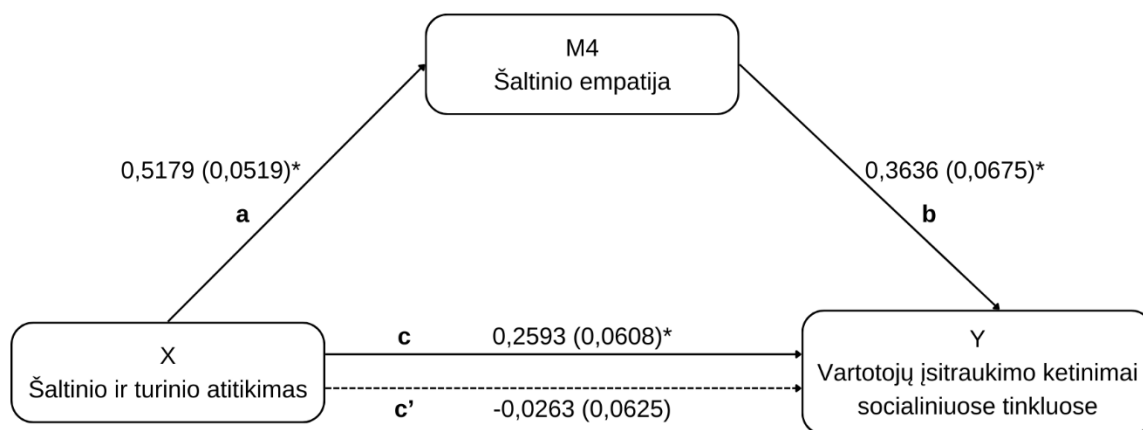
R	Priklausomas kintamasis											
	M4: Šaltinio empatija				Y: Įsitraukimo ketinimai				Y: Įsitraukimo ketinimai			
	→	Koef.	SE	p	→	Koef.	SE	p	→	Koef.	SE	p
Konstanta	iM →	2,2537	0,2327	0,0000	iY →	0,5751	0,2866	0,0455	iY →	1,5064	0,2729	0,0000
X: Šaltinio ir turinio atitikimas	a →	0,5179	0,0519	0,0000	c'→	-0,0263	0,0625	0,6743	c →	0,2593	0,0608	0,0000
M4: Šaltinio empatija	-	-	-	-	b→	0,3636	0,0675	0,0000	-	-	-	-
Modelis	$R^2 = 0,2124$, $F(1;370) = 99,7608$, $p=0,0000$				$R^2 = 0,2872$, $F(6;365) = 24,5143$, $p=0,0000$				$R^2 = 0,0468$, $F(1;370) = 18,1735$, $p=0,0000$			

49 lentelėje pateikti rezultatai patvirtina reikšmingą netiesioginį poveikį per šaltinio empatiją: 95 % pasikliautinis intervalas nuo 0,1050 iki 0,2882 neapima nulio, todėl galima tvirtinti, kad šaltinio empatija veikia kaip mediatorius tarp šaltinio ir turinio atitikimo bei vartotojų įsitraukimo ketinimų.

49 lentelė. Tiesioginė, netiesioginė (per šaltinio empatiją) ir suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams (N=372)

Kelias	Efektas (EF)	95% pasikliautinis intervalas	
		LLCI	ULCI
Tiesioginė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → įsitraukimo ketinimai (c’)	-0,0263	-0,1492	0,0966
Netiesioginė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → šaltinio empatija → įsitraukimo ketinimai (a*b)	0,1883	0,1050	0,2882
Suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → įsitraukimo ketinimai (c)	0,2593	0,1397	0,3788

Ši sąsaja vizualiai pavaizduota 9 paveiksle.



9 pav. Šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams per šaltinio empatiją

* $p < 0,001$

Apibendrinant galima teigti, kad hipotezė **H8c**, teigianti, kad „šaltinio empatija teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose“, yra pagrįsta empiriniais duomenimis ir patvirtinama.

Remiantis pateiktais duomenimis (žr. 50 lentelėje) galima teigti, kad egzistuoja statistiškai reikšminga šaltinio ir turinio atitikimo įtaka šaltinio palankumui (a kelias; $p = 0,000$) ir įsitraukimo ketinimams (b kelias; $p = 0,2633$). Šaltinio ir turinio atitikimas suminio efekto įtaka įsitraukimo ketinimams yra taip pat statistiškai reikšminga (c kelias; $p = 0,000$).

50 lentelė. Regresijos modelių, tiriant ryšį tarp šaltinio ir turinio atitikimo, pasitikėjimo šaltiniu ir įsitraukimo ketinimų, rezultatai

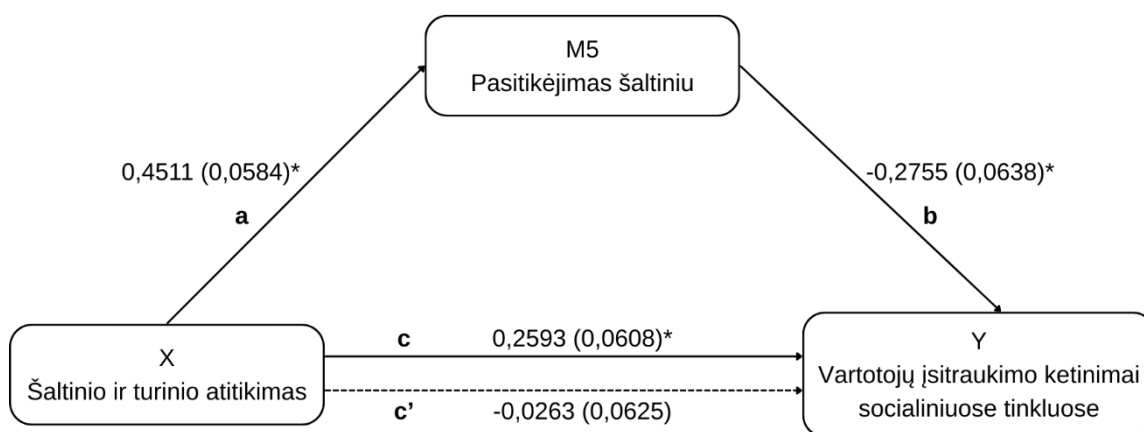
R	Priklausomas kintamasis											
	M5: Pasitikėjimas šaltiniu				Y: Įsitraukimo ketinimai				Y: Įsitraukimo ketinimai			
	→	Koef.	SE	p	→	Koef.	SE	p	→	Koef.	SE	p
Konstanta	iM →	2,7053	0,2619	0,0000	iY →	0,5751	0,2866	0,0455	iY →	1,5064	0,2729	0,0000
X: Šaltinio ir turinio atitikimas	a →	0,4511	0,0584	0,0000	c'→	-0,0263	0,0625	0,6743	c →	0,2593	0,0608	0,0000
M5: Pasitikėjimas šaltiniu	-	-	-	-	b→	-0,2755	0,0638	0,0000	-	-	-	-
Modelis	$R^2 = 0,1390$, $F(1;370) = 59,7346$, $p = 0,0000$				$R^2 = 0,2872$, $F(6;365) = 24,5143$, $p = 0,0000$				$R^2 = 0,0468$, $F(1;370) = 18,1735$, $p = 0,0000$			

51 lentelėje pateikiama tiesioginė, netiesioginė ir suminė šaltinio ir turinio atitikimo įtaka įsitraukimo ketinimams.

51 lentelė. Tiesioginė, netiesioginė (per šaltinio empatiją) ir suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams (N=372)

Kelias	Efektas (EF)	95% pasikliautinis intervalas	
		LLCI	ULCI
Tiesioginė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → įsitraukimo ketinimai (c’)	-0,0263	-0,1492	0,0966
Netiesioginė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → pasitikėjimas šaltiniu → įsitraukimo ketinimai (a*b)	-0,1243	-0,2136	-0,0586
Suminė šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams			
Šaltinio ir turinio atitikimas → įsitraukimo ketinimai (c)	0,2593	0,1397	0,3788

Iš rezultatų, pateiktų 51 lentelėje, matoma, kad atitikimas didina pasitikėjimą šaltiniu, kuris savo ruožtu mažina įsitraukimą, todėl netiesioginio ryšio įtaka nėra patvirtinama – tą rodo netiesioginio ryšio neigiamas koeficientas. Atliktos analizės rezultatai pateikiami 10 paveiksle.



10 pav. Šaltinio ir turinio atitikimo įtaka vartotojų įsitraukimo ketinimams per pasitikėjimą šaltiniu

* $p < 0,001$

Apibendrinant, analizės rezultatai rodo, šaltinio ir turinio atitikimas turi reikšmingą teigiamą poveikį turinio patikimumui ($\beta = 0,2967$; $p = 0,0000$), kuris savo ruožtu padidina vartotojų ketinimus įsitraukti į turinį socialiniuose tinkluose (EF = 0,1655; [0,0920; 0,2565]). Be to, šaltinio palankumas taip pat medijuoja šaltinio ir turinio atitikimo ryšį su vartotojų įsitraukimo ketinimais (EF = 0,0960; [0,0495; 0,1541], $p = 0,000$). Dar viena reikšminga medijuojanti dimensija yra šaltinio empatija ($\beta = 0,3636$; $p = 0,0000$) – ji teigiamai medijuoja šaltinio ir turinio atitikimo poveikį įsitraukimui (EF = 0,1883; [0,1050; 0,2882]). Tai rodo, kad kai vartotojai suvokia šaltinį kaip empatišką, jie dažniau įsitraukia į turinį, nes tokiu būdu stiprinamas pasitikėjimas šaltiniu ir turiniu.

Tuo tarpu šaltinio kompetencija ir pasitikėjimas nemedijuoja šaltinio ir turinio atitikimo poveikio įsitraukimui. Šaltinio kompetencijos koeficientas ($\beta = -0,0814$; $p = 0,2633$) rodo, kad kompetencija neturi reikšmingo poveikio vartotojo įsitraukimo ketinimams, nes šis kintamasis nesukelia reikšmingų pokyčių šaltinio ir turinio atitikimo įtakos įsitraukimui. Be to, pasitikėjimas šaltiniu turi neigiamą koeficientą ($\beta = -0,2755$; $p = 0,0000$).

52 lentelėje pateikta hipotezių tikrinimo rezultatų suvestinė, kurioje apibendrinami visų tirtų ryšių tarp suvokiamo šaltinio ir turinio atitikimo bei kitų tyrimo kintamųjų rezultatai, nurodant, kurios hipotezės buvo patvirtintos remiantis atliktomis statistinėmis analizėmis, o kurios – nepatvirtino dėl silpno poveikio dydžio ar kitų priežasčių.

52 lentelė. Hipotezių tikrinimo rezultatų suvestinė (N = 372)

Hipotezė	Rezultatas	Pagrindimas
H1: DI sukurtas ir pažymėtas rinkodaros turinys (palyginti su nepažymėtu) sukels didesnį suvokiamą turinio ir jo šaltinio atitikimą.	Patvirtinta	Žr. 24 lentelėje
H2: DI sukurtas ir pažymėtas rinkodaros turinys (palyginti su nepažymėtu) sukels didesnį suvokiamą turinio ir jo šaltinio atitikimą kai turinys turi sąsają su DI.	Patvirtinta	Žr. 25–27 lentelėse
H3a: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų suvokiamą šaltinio kompetenciją.	Patvirtinta	Žr. 28–35 lentelėse
H3b: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų suvokiamą šaltinio empatiją.	Patvirtinta	
H3c: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų suvokiamą šaltinio palankumą.	Patvirtinta	
H3d: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų pasitikėjimą šaltiniu.	Patvirtinta	
H4: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų suvokiamą turinio patikimumą.	Patvirtinta	Žr. 36–37 lentelėse
H5: Vartotojų suvokiamas turinio ir jo šaltinio atitikimas teigiamai veikia vartotojų įsitraukimo ketinimus socialiniuose tinkluose.	Nepatvirtinta	Žr. 38–39 lentelėse
H6a: Vartotojų suvokiama šaltinio kompetencija teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;	Nepatvirtinta	Žr. 40–41 lentelėse
H6b: Vartotojų suvokiama šaltinio empatija teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;	Patvirtinta	
H6c: Vartotojų suvokiamas šaltinio palankumas teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;	Patvirtinta	
H6d: Vartotojų suvokiamas pasitikėjimas šaltiniu teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose;	Nepatvirtinta	
H7: Vartotojų suvokiamas turinio patikimumas teigiamai veikia jų įsitraukimo ketinimus socialiniuose tinkluose.	Patvirtinta	Žr. 44–51 lentelėse
H8a: Šaltinio kompetencija teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose;	Nepatvirtinta	
H8b: Šaltinio empatija teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose;	Patvirtinta	
H8c: Šaltinio palankumas teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose;	Patvirtinta	
H8d: Pasitikėjimas šaltiniu teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose.	Patvirtinta	
H8: Turinio patikimumas teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose.	Patvirtinta	Žr. 42–43 lentelėse
H9: Šaltinio patikimumas teigiamai medijuoja ryšį tarp šaltinio ir turinio atitikimo ir vartotojų įsitraukimo socialiniuose tinkluose	Patvirtinta	

Nors hipotezė H5, H6a, H6d, H8a buvo atmestos, visos kitos hipotezės (H1–H4, H6b, H6c, H7, H8b–H8d, H9) buvo patvirtintos. Kaip rodo analizės rezultatai, DI autorystės žymėjimas reikšmingai padidino šaltinio ir turinio atitikimo suvokimą, ypač kai turinys buvo susijęs su DI tema. Savo ruožtu,

šis suvoktas atitikimas turėjo teigiamą poveikį šaltinio patikimumo dimensijoms – kompetencijai, palankumui, empatijai ir pasitikėjimui – bei turinio patikimumo vertinimui. Taip pat nustatyta, kad turinio patikimumas reikšmingai didino vartotojų išitraukimo ketinimus, tačiau pats atitikimas ar kai kurios šaltinio savybės (pavyzdžiui, empatija) šio ketinimo tiesiogiai nepaveikė. Tai rodo, jog išitraukimą lemia labiau turinio patikimumas nei pavieniai šaltinio ar atitikimo aspektai. Gauti rezultatai patvirtina, kad kontekstualiai pagrįstas DI žymėjimas gali sustiprinti vartotojų pasitikėjimą tiek turiniu, tiek šaltiniu.

4.6. Dirbtinio intelekto sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose tyrimo rezultatų diskusija

Šio baigiamojo magistro projekto tikslas buvo ištirti, kaip DI sukurto rinkodaros turinys ir jo autorystės žymėjimas veikia vartotojų suvokiamą šaltinio ir turinio atitikimą, suvokiamą patikimumą bei išitraukimo ketinimus socialiniuose tinkluose. Atlikta mokslinės literatūros analizė parodė, kad tyrimų rezultatai šioje srityje nėra vienareikšmiai – mokslinėje bendruomenėje vyrauja prieštaringos nuomonės apie vartotojų požiūrį į DI sukurtą rinkodaros turinį socialiniuose tinkluose. Iki šiol atlikti tyrimai rodo, kad vartotojų reakcijos į įvairias turinio žymas – tokias kaip „remiama“, „klaidinga“ ar „užginčytas turinys“ – dažniausiai priklauso nuo pačios žymos tipo, jos pateikimo formos ir vartotojų išankstinio supratimo apie žymos reikšmę (Boerman et al., 2017; De Veirman & Hudders, 2020; Clayton ir kt., 2020; Altay & Gilardi, 2024). Tokios žymos dažnai sukelia vartotojų skepticizmą, mažina pasitikėjimą turiniu ir mažina išitraukimo tikimybę, nes atskleidžia galimą komercinį ar manipuliatyvų turinio pobūdį.

Panašios tendencijos stebimos ir naujausiuose tyrimuose, nagrinėjančiuose DI sukurtą turinį. Literatūros analizė atskleidė, kad autorystės žymėjimas etikete „sukurta DI“ vartotojų gali kelianti abejones dėl turinio patikimumo, ypač jei vartotojai turi ribotą supratimą apie DI technologijų veikimo principus (Fisher et al., 2024; Altay & Gilardi, 2024; Lim & Schmalzle, 2024). Nors kai kurie tyrimai (Kirkby et al., 2023, Park et al., 2024) fiksuoja neutralų ar net teigiamą vartotojų požiūrį į DI autorystės turinį, dauguma šaltinių atskleidžia, kad DI turinio autorystės žymėjimas mažina pasitikėjimą tokiu turiniu, išitraukimą į jį, net jei jo objektyvi kokybė nėra prastesnė nei žmogaus sukurto (Zhang & Gosline, 2023; Carney et al., 2024; Gu et al., 2024; Lee et al., 2025).

Nepaisant šių tyrimų rezultatų, vis dar lieka neatsakyta, ar DI autorystės žymėjimas turi vienodą poveikį skirtingiems turinio tipams. Dauguma iki šiol atliktų tyrimų daugiausia dėmesio skyrė bendro pobūdžio rinkodaros ar informaciniam turiniui, kuris tiesiogiai nesusijęs su pačia DI tematika. Dėl to išryškėja teorinė spraga – lieka neaišku, ar DI žyma vartotojų yra suvokiama skirtingai, kai turinys susijęs su DI produktais, palyginti su atvejais, kai turinys su DI nėra susijęs. Be to, nepakankamai ištirta, kaip DI autorystės žymos socialiniuose tinkluose veikia vartotojų reakcijas į mišraus pobūdžio turinį, kuriame derinami tiek tekstiniai, tiek vizualiniai elementai – būtent toks formatas yra būdingas rinkodaros komunikacijai socialiniuose tinkluose. Šiame baigiamajame magistro darbe nagrinėti reiškiniai – DI autorystės žymėjimas, šaltinio ir turinio atitikimas, suvokiamas patikimumas bei vartotojų išitraukimo ketinimai – iki šiol nebuvo sistemingai tyrinėti viename tyrime kaip tarpusavyje susiję veiksniai. Todėl šiuo projektu buvo siekiama ne tik prisidėti prie esamos teorinės bazės papildymo, bet ir gilinti supratimą apie tai, kaip DI turinio žymėjimas bei vartotojų suvokiamos šaltinio ir turinio savybės veikia pasitikėjimą turiniu ir išitraukimą į jį šiuolaikinėse socialinių tinklų platformose.

Remiantis mokslinės literatūros analize buvo keliama prielaida, kad DI sukurtų įrašų žymėjimas gali turėti įtakos vartotojų suvokiamam turinio ir jo šaltinio atitikimui. Ši prielaida grįsta schemos atitikimo teorija (Mandler, 1982; Meyers-Levy & Tybout, 1989), teigiančia, kad žmonės informaciją vertina palankiau, kai ji atitinka jų turimas kognityvines schemas. Suvokiamas šaltinio ir turinio atitikimas šiame projekte buvo suprantamas kaip vartotojo subjektyvus vertinimas, kiek reklamos žinutė logiškai dera su ją pateikusių šaltiniu. Neatitikimas tarp turinio ir jo šaltinio – pavyzdžiui, kai reklaminę žinutę, susijusią su žmogaus kompetencijomis, pateikia DI – gali sukelti kognityvinį diskomfortą ar mažinti pasitikėjimą turiniu (Peracchio & Tybout, 1996; Yoon, 2013). Šią teorinę prielaidą sustiprina empiriniai tyrimai, rodantys, kad pažymėtas DI autorystės turinys, ypač kai nėra susijęs su DI tematika, dažnai vertinamas skeptiškiau nei nepažymėtas arba žmogaus kurtas turinys (Zhang & Gosline, 2023; Lee et al., 2025). Tuo tarpu kai turinio tematika yra susijusi su pačiu DI – pavyzdžiui, kai reklamuojamas DI paremtas sprendimas – žymėjimas „sukurta DI“ gali būti suvokiamas kaip natūralus ir mažiau trikdomas, nes tarp turinio ir jo šaltinio atsiranda loginė sąsaja. Taip pat literatūros analizės rezultatai rodo, jog šaltinio ir turinio atitikimas gali daryti poveikį vartotojų suvokiamam patikimumui bei įsitraukimo ketinimams socialiniuose tinkluose. Suvokiamas šaltinio patikimumas šiame projekte buvo suprantamas pagal McCroskey'aus ir Teven'o (1999) modelį, kuris apima tris tarpusavyje susijusias dimensijas – kompetenciją, palankumą ir pasitikėjimo vertumą. Vis dėlto atlikus tiriamąją faktoriinę analizę su šaltinio patikimumo teiginiais paaiškėjo, kad konstrukto struktūra yra labiau fragmentuota nei teoriškai numatyta: respondentai šiuos požymius vertino kaip atskirus, o šaltinio patikimumas empiriškai išsiskaidė į keturias atskiras dimensijas – kompetenciją, empatiją, palankumą ir pasitikėjimą. Tokiu būdu teorinė skalės samprata buvo patikslinta remiantis empiriniais rezultatais. Tuo tarpu suvokiamas turinio patikimumas siejamas su turinio aiškumu, tikslumu ir objektyvumu, o įsitraukimo ketinimai – su vartotojo noru reaguoti į turinį socialinių tinklų platformose: komentuoti, dalytis ar spausti „patinka“. Kai šaltinis ir turinys yra suvokiami kaip tarpusavyje atitinkantys, turinys gali būti laikomas įtikinamesniu, o šaltinis – labiau kompetentingu (Hovland et al., 1953; Appelman & Sundar, 2016; Ecker & Antonio, 2021). Toks suvokiamas atitikimas gali lemti palankesnę vartotojų atsaką, įskaitant ketinimą įsitraukti į turinį socialiniuose tinkluose – spausti mygtuką „patinka“, dalytis ar komentuoti (Onofrei et al., 2022). Remiantis teorinėmis nuostatomis ir mokslinės literatūros analize, šiame tyrime buvo iškeltos 8 hipotezės, kurios vėliau buvo detalizuotos bei tikrinamos pasitelkiant kiekybinio tyrimo metodą – internetu vykdytą eksperimentą. Tyrimo instrumente buvo įtrauktas atviras klausimas, kuris leido atlikti tiek kiekybinę, tiek kokybinę turinio analizę, suteikiant galimybę gauti papildomų vertingų įžvalgų.

Tyrimo rezultatai rodo, jog DI autorystės žyma turėjo statistiškai reikšmingą poveikį respondentų šaltinio ir turinio atitikimo suvokimui. Nepriklausomų imčių t-testas parodė, kad vidurkių skirtumas buvo reikšmingas ($p < 0,001$), o efektas – vidutinio stiprumo (Cohen $d \approx 0,38$). Tai reiškia, kad socialiniame tinkle „Facebook“ naudojama žyma „AI info“ DI turiniui žymėti veikia kaip orientyras, stiprinantis vartotojo išpūdį apie komunikacijos vientisumą ir patikimumą. Be to, dvifaktorių dispersijos analize nustatyta, kad toks DI turinio žymėjimas ne tik sustiprino šaltinio ir turinio atitikimo suvokimą, bet ir sąveikavo su turinio tematika – nustatyta reikšminga sąveika tarp žymėjimo ir turinio tematikos ($p = 0,030$). Šie rezultatai atliepia schemos atitikimo teoriją (Choi, 2024), kuri teigia, kad vartotojai lengviau supranta ir priima informaciją, kai ji atitinka jų lūkesčius apie šaltinį, temą ir pateikimo kontekstą. Šiuo atveju, kai turinys apie inovaciją, turinčią aiškią sąsają su DI, yra sukurtas DI ir pažymėtas kaip DI sukurtas, vartotojai pastebi logišką sąsają tarp jų – tema, kilmė ir

žyma vienas kitą sustiprina. Dėl šios dermės turinys atrodo nuoseklus ir nekeliantis kognityvinio disonanso, todėl jį lengviau suprasti ir priimti.

Sentimentų analizė atskleidė, kad vartotojų emocinis atsakas į manipuliatyvų turinį daugeliu atvejų buvo neutralus. Vis dėlto reikšminga dalis atsakymų pasižymėjo neigiamu tonu – daugiausiai kritikos sulaukė kalbos stilius, sakinių logika ir vizualinių elementų suderinamumo stoka. Tai rodo, kad net ir nepažymėtas turinys, turintis vizualių neatitikimų su realiais vaizdais ar pasižymintis neįprasta sakinių struktūra, gramatinėmis ar skyrybos klaidomis, gali būti priskirtas DI. Teigiami atsakymai dažniausiai pasireiškė tais atvejais, kai turinys buvo pateiktas aiškiai, be gramatinių klaidų ir su nuoseklia žinute. Tokie vertinimai pabrėžia, kad DI komunikacijos sėkmė priklauso ne tik nuo technologinio sprendimo, bet ir nuo jo stilistinės realizacijos.

Tyrimo duomenys patvirtina, kad didesnis suvokiamas atitikimas tarp šaltinio ir turinio reikšmingai padidina abiejų šaltinio ir turinio patikimumo vertinimus, o tai sutampa su Metzger'io et al. (2003) siūloma daugiamate patikimumo samprata. Regresinių analizių kiekvienai šaltinio patikimumo dimensijai rezultatai patvirtino statistškai reikšmingą teigiamą ryšį tarp šaltinio ir turinio atitikimo suvokimo bei šaltinio patikimumo – empatijos, palankumo, pasitikėjimo ir kompetencijos. Šis nuoseklus ryšys rodo, kad kuo aiškiau vartotojai identifikuoja atitikimą tarp skelbiamo turinio ir jį pateikusios organizacijos ar asmens, tuo aukštesniais balais vertina šaltinio patikimumą. Svarbu pastebėti, kad stipriausias ryšys nustatytas tarp atitikimo suvokimo ir šaltinio empatijos vertinimo, o tai leidžia teigti, kad nuoseklus turinys padeda suformuoti ne tik racionalų, bet ir emocinį pasitikėjimo pagrindą. Tai ypač svarbu komunikacijos srityje, kur vartotojo įsitraukimas ir pasitikėjimas neretai priklauso nuo to, kiek turinį publikuojanti organizacija atrodo supratinga ir nesavanaudiška. Regresinės analizės rezultatai parodė, kad šaltinio ir turinio atitikimo suvokimas taip pat reikšmingai prognozuoja paties turinio patikimumo įvertinimą. Tai reiškia, kad turinys, kuris auditorijai atrodo „tinkantis“ šaltiniui, laikomas ir pačiuose turinio lygmenyse labiau patikimu. Tokia sąsaja leidžia manyti, kad vartotojai nevertina turinio izoliuotai, bet remiasi kontekstiniais signalais – jei turinys „logiškai dera“ prie šaltinio tapatybės, jį lengviau priimti kaip patikimą.

Priešingai nei tikėtasi, tyrimo rezultatai parodė, jog vartotojų suvokiamas turinio ir jo šaltinio atitikimas nedarė įtakos vartotojų įsitraukimo ketinimams socialiniuose tinkluose. Nors ryšys tarp šaltinio ir turinio atitikimo suvokimo bei įsitraukimo ketinimų buvo statistškai reikšmingas, efektas buvo labai silpnas, o tai rodo, kad praktinė reikšmė yra minimali. Todėl galima teigti, jog suvoktas atitikimas tarp šaltinio ir turinio nėra pakankamas faktorius, skatinantis auditoriją reaguoti į turinį elgseniniu būdu – komentuoti, dalintis ar spausti „patinka“. Tai leidžia daryti išvadą, kad kiti veiksniai turi svarbesnį vaidmenį formuojant vartotojų elgesį socialiniuose tinkluose, pavyzdžiui – šaltinio ir turinio patikimumas, kuris tiesiogiai veikia vartotojų įsitraukimo ketinimus. Šis ryšys yra svarbus socialinių tinklų kontekste, kuriame, kaip rodo Onofrei'is et al. (2022) ir Boerman'as et al. (2017), pasitikėjimas žinute bei jos autoriumi dažnai tampa pagrindiniu motyvu įsitraukti – spausti „patinka“, komentuoti ar dalintis turiniu. Turinio patikimumas buvo reikšmingas ir gana stiprus įsitraukimo ketinimų prediktorius, o iš keturių tirtų šaltinio patikimumo dimensijų teigiamą poveikį įsitraukimo ketinimams turėjo empatija ir palankumas.

Tai leidžia patvirtinti ankstesnius tyrimus, kurie rodo, kad ne tik turinio kokybė, bet ir jo suderinamumas su auditorijos lūkesčiais bei šaltinio empatiškumas ir palankumas vartotojo atžvilgiu lemia elgseninį vartotojų atsaką. Įdomu tai, jog šaltinio pasitikėjimo dimensija turėjo neigiamą poveikį vartotojų įsitraukimo ketinimams. Šis netikėtas rezultatas leidžia daryti prielaidą, kad

vartotojai sąveikauja su turiniu ne tiek dėl racionalaus šaltinio įvertinimo (pavyzdžiui, kompetencijos), bet labiau dėl emocinių impulsų – kai šaltinis atrodo šiltas, palankus ir empatiškas. Tuo tarpu aukštas pasitikėjimo lygis gali sukelti priešingą efektą: jei šaltinis laikomas visiškai patikimu, vartotojai gali nebejausti poreikio papildomai reaguoti į turinį. Galima manyti, kad tokiais atvejais išitraukimo motyvacija sumenksta ne dėl nepasitikėjimo, o dėl sumažėjusio stimulo – informacinis, patikimas turinys gali būti suvokiamas kaip akivaizdus, nekeliantis abejonių, todėl nereikalaujantis papildomos vartotojo reakcijos. Kita vertus, socialiniuose tinkluose vartotojai dažnai išitraukia į turinį, kuris jiems atrodo provokuojantis, prieštaringas ar neatitinkantis jų lūkesčių. Tad racionalus, informatyvus turinys, net jei jis patikimas, ne visada veikia kaip išitraukimą skatinantis stimulus – veikiau priešingai, jis gali būti priimtas pasyviai.

Tyrimo apribojimai:

1. Tyrime taikyti eksperimentiniai scenarijai buvo grindžiami informacinio pobūdžio turiniu, kuris nėra skirtas tiesioginiam prekių ar paslaugų pardavimų skatinimui. Nei vienas iš pateiktų įrašų nesiūlė konkretaus komercinio veiksmo, todėl tyrimo rezultatai yra labiau tinkami vertinant nekomercinio pobūdžio komunikaciją (pavyzdžiui, organizacijų ar viešųjų įstaigų komunikaciją), o jų taikymas tiesioginiams komercinės reklamos atvejams (pavyzdžiui, prekių ženklų kampanijoms) gali būti ribotas. Dėl šios priežasties tyrimo išvados negali būti tiesiogiai perkeltamos į verslo įmonių rinkodaros praktiką, ypač tuomet, kai siekiama skatinti vartotojų pirkimo elgseną.
2. Tyrime naudotas eksperimentinis scenarijus rėmėsi „Facebook“ platformos simuliacija, kurioje vartotojui pirmiausia pateikiama DI autorystės žyma, o tik tuomet tekstas bei vizualas. Tokia turinio struktūra gali veikti kitaip nei kitose platformose, tokiose kaip „Instagram“ ar „TikTok“, kur dominuoja vizualinė medžiaga, o tekstas ir žymos dažnai yra rodomi įrašo pabaigoje ar išvis paslepiami. Dėl šios priežasties tyrimo rezultatų taikymas kitų platformų turiniui gali būti ribotas.
3. Tyrimo konstrukto – šaltinio pasitikėjimo – matavimui buvo pasitelkta skalė, kurios struktūra buvo išgryninta atliekant faktorinę analizę. Šios analizės metu paaiškėjo, kad konstrukto dimensijos išsiskyrė kitaip nei teoriškai tikėtasi, o tai riboja galimybes rezultatus tiesiogiai lyginti su kitais tyrimais, kuriuose buvo taikyta atitinkama skalė.
4. Šiame projekte vartotojų išitraukimo ketinimai buvo vertinami remiantis jų subjektyviais atsakymais anketose. Tai reiškia, kad tyrime analizuotas elgesys buvo deklaratyvus, o ne objektyviai stebėtas. Atsižvelgiant į tai, kad reali elgsena socialiniuose tinkluose (komentavimas, pasidalinimas ar „patinka“ paspaudimai) ne visuomet atitinka apklausose išreikštus ketinimus, tyrimo rezultatai gali neatspindėti faktinės sąveikos su turiniu.

Tolimesnės tyrimų kryptys:

1. Ateities tyrimuose tikslinga taikyti reprezentatyvią tikimybinę atranką, leidžiančią patikimiau apibendrinti rezultatus visai šalies populiacijai. Tai ypač svarbu, siekiant teikti praktines rekomendacijas Lietuvos rinkodaros komunikacijos praktikams.
2. Tolesniuose tyrimuose vertėtų išplėsti tiriamų temų spektrą, siekiant nustatyti, ar sąsajos su DI stebimi efektai pasikartoja ir kitose inovacijų tematikose. Būtų tikslinga į tyrimą įtraukti skirtingų sričių inovacijas – neapsiribojant sveikata ir maistu, bet analizuojant ir, pavyzdžiui, technologijų, energetikos, transporto ar švietimo sritis. Tai leistų įvertinti, ar DI žymėjimo bei šaltinio ir turinio atitikimo poveikis yra bendro pobūdžio, ar specifinis tam tikroms temoms.
3. Ateities tyrimuose rekomenduojama analizuoti skirtingų socialinių platformų turinį, tokių kaip „Instagram“ ar „TikTok“, kurios pasižymi skirtingomis turinio struktūromis. Tai leistų geriau

- suprasti, kaip vizualinė medžiaga – nuotraukos ar vaizdo įrašai – ir tekstas sąveikauja su vartotojais įvairiose platformose bei kaip šios struktūros veikia vartotojų reakcijas ir įsitraukimą.
4. Būtų prasminga atlikti ilgalaikius tyrimus, siekiant įvertinti, ar vartotojų požiūris į DI sukurtą turinį keičiasi laikui bėgant. Ilgalaikis stebėjimas galėtų padėti suprasti, ar įpročiai ir vertinimai formuojasi palaipsniui, ar išlieka stabilūs.
 5. Tiriant vartotojų įsitraukimo elgseną rekomenduojama integruoti kiekybinius rodiklius, tokius kaip komentarų skaičius, pasidalijimų dažnumą ir „patinka“ paspaudimų rodiklius, siekiant įvertinti vartotojų įsitraukimą realioje socialinių tinklų aplinkoje. Objektyvūs vartotojų elgsenos duomenys, surinkti natūralioje aplinkoje, leistų geriau įvertinti faktinę jų sąveiką su turiniu ir pateikti išsamesnes įžvalgas apie vartotojų įsitraukimą.
 6. Rekomenduojama tirti asmeninius kintamuosius, pavyzdžiui, vartotojų požiūrį į technologijas, informacinį raštingumą ar autentiškumo poreikį. Tokie kintamieji galėtų paaiškinti skirtingą vartotojų jautrumą DI žymėjimui bei šaltinio ir turinio patikimumui.

Išvados ir rekomendacijos

Teoriškai ir empiriškai išanalizavus DI sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikį vartotojų atsakui socialiniuose tinkluose, daromos šios išvados:

1. DI sukurto rinkodaros turinio ir jo autorystės žymėjimo poveikio vartotojų atsakui socialiniuose tinkluose tyrimas yra aktualus dėl prieštaringų vartotojų vertinimų, skirtingų turinio kontekstų bei reikšmingų tyrimų spragų. Remiantis įvairių mokslinių tyrimų apžvalga, nagrinėjančia vartotojų požiūrį į DI valdomų sprendimų sukurtą rinkodaros turinį, nustatyta, kad DI taikymas komunikacijoje sukelia nevienareikšmį vartotojų atsaką. Vieni tyrimai rodo, kad vartotojai DI sukurtą turinį vertina taip pat kaip ir žmogaus kurtą, o DI autorystės žymėjimas neturi reikšmingo poveikio jų požiūriui (Kirkby et al., 2023). Tačiau kita tyrimų grupė atskleidžia, kad aiškus DI autorystės žymėjimas dažnu atveju sumažina teigiamą turinio vertinimą – ypač tais atvejais, kai vartotojai suvokia DI kaip mažiau autentišką ar patikimą turinio šaltinį (Zhang & Gosline, 2023; Lim & Schmalzle, 2024; Shantatula et al., 2024; Lee et al., 2025). Svarbu pažymėti, kad dauguma esamų tyrimų orientavosi į bendrinį rinkodaros turinį, neatsižvelgiant į jo tematinę ryšį su DI. Dėl to trūksta įžvalgų apie tai, kaip kinta vartotojų požiūris, kai DI sukurtas turinys yra susijęs su inovatyviais DI taikymo atvejais – pavyzdžiui, kai publikuojamas turinys yra apie DI pagrindu sukurtus produktus. Nežinoma, ar tokio turinio ir jo autorystės žymėjimo sąsaja yra vartotojų priimama kaip prasminga ir sustiprinanti pasitikėjimą, ar atvirkščiai – kelia papildomų abejonių dėl turinio manipuliatyvumo. Be to, eksperimentiniai tyrimai dažniausiai apsiribojo vieno tipo turiniu (dažniausiai tekstiniu), ignoruojant socialinių tinklų specifiką, kurioje tekstinis turinys įprastai publikuojamas kartu su vizualiniu. Kadangi DI autorystės žymėjimas vis dar yra palyginti naujas reiškinys, lieka neaišku, ar vartotojai jį taikantys visam turinio rinkiniui, ar tik jo dalims. Dėl to šio darbo kontekste analizuojama tyrimų problema – kaip DI sukurtas turinys ir jo žymėjimas veikia vartotojų atsaką socialiniuose tinkluose – laikytina aktualia tiek teoriniu, tiek praktiniu požiūriu. Ji atveria galimybes užpildyti žinių spragas apie DI komunikacijos priimtinumą bei prisidėti prie įrodymais grįstų rekomendacijų formavimo rinkodaros specialistams, ypač kuriant su mokslu ar technologijomis susijusį DI turinį.
2. Atlikta išsami literatūros analizė atskleidė teorines prieigas, leidžiančias paaiškinti vartotojų elgseną susidūrus su DI sukurtu ir pažymėtu turiniu. Remiantis algoritmų vengimo teorija (Dietvorst et al., 2015; Castelo et al., 2019), vartotojai linkę skeptiškai vertinti DI sprendimus, ypač kai šie sprendimai siejami su subjektyvumu, kūrybiškumu ar emocine reikšme – tais aspektais, kurie tradiciškai laikomi žmogaus kompetencijos sritimi. Taip pat aktualus įtaigos atpažinimo modelis, kuris teigia, jog vartotojai, atpažinę reklamos šaltinio ar technologijos manipuliatyvias intencijas, tampa labiau linkę gintis – tai pasireiškia didesniu skepticizmu, mažesniu pasitikėjimu turiniu ir silpnesniu įsitraukimu (Friestad & Wright, 1994; Boerman et al., 2017; Clayton et al., 2020). Atsižvelgiant į šiuos teorinius pagrindus, darbe suformuotas konceptualus modelis, kuriame turinio autorystės žymėjimas (pažymėta DI ar nepažymėta) laikomas pagrindiniu nepriklausomu kintamuoju. Modelyje taip pat įtraukti moderuojantys ir medijuojantys ryšiai – turinio sąsaja su DI veikia kaip moderuojantis veiksnys, o suvokiamas šaltinio ir turinio atitikimas – kaip mediatorius, paaiškinantis, kaip žymėjimas paveikia vartotojų atsaką. Taip pat, modelyje analizuotas ir šaltinio bei turinio patikimumas kaip galimas mediacinis mechanizmas, aiškinantis kaip šaltinio ir turinio atitikimas veikia vartotojų įsitraukimo ketinimus. Modelis grindžiamas prielaida, jog pažymėtas DI turinys gali sustiprinti suvokimą apie šaltinio ir turinio suderinamumą, jei turinys yra susijęs su DI tematika. Tokia sąveika aiškinama schemos atitikimo teorija, kuri teigia, kad jei turinys atitinka vartotojo išankstinius lūkesčius apie jo šaltinį

(pavyzdžiui, kad DI sukurtas turinys pristato DI produktus), jis laikomas nuoseklesniu, aiškesniu ir todėl patikimesniu. Priešingai, neatitikimas tarp šaltinio ir turinio – pavyzdžiui, kai DI kurtas turinys pristato visiškai su DI nesusijusį produktą – gali kelti abejones dėl autentiškumo, paskatinti nepasitikėjimą ir sumažinti vartotojų įsitraukimą. Galiausiai, literatūros analizė parodė, kad DI žymėjimas, ypač kai nesutampa su vartotojo išankstinėmis nuostatomis ar lūkesčiais, gali sukelti „neatitikimo“ efektą, mažinantį patikimumą ir įsitraukimą – ypač, jei DI naudojamas tematiname kontekste, kuriame svarbios žmogiškos savybės (Lee et al., 2025; Zhang & Gosline, 2023). Todėl šio darbo konceptualus modelis, apimantis autorystės žymėjimą, turinio teminį kontekstą bei vartotojų suvokiamą šaltinio ir turinio atitikimą, yra teoriškai pagrįstas ir leidžia patikrinti sudėtingus vartotojų elgsenos dėsningumus, susijusius su DI naudojimu rinkodaroje socialiniuose tinkluose.

3. Tyrimo metodologija buvo pagrįsta sistemingu eksperimentiniu dizainu, leidžiančiu patikimai įvertinti DI sukurtą turinio autorystės žymėjimą ir tematinio konteksto poveikį vartotojų atsakai socialiniuose tinkluose. Siekiant empiriškai ištirti, kaip DI turinio autorystės žymėjimas ir turinio sąsaja su DI tematika veikia vartotojų atsaką, buvo pasirinktas kiekybinis 2x2 tarpgrupinis eksperimentinis dizainas, pagrįstas keturiais skirtingais scenarijais. Šis dizainas imitavo realistiškas socialinių tinklų žinutes, kuriose manipuliuojama dviem nepriklausomais kintamaisiais: turinio autorystės žymėjimu (pažymėtas ar ne) ir turinio tematinio ryšiu su DI (yra ar nėra). Tokiu būdu buvo užtikrinta galimybė sistemiškai tirti priežastinius ryšius tarp manipuliuojamų veiksnių ir vartotojų atsako formų. Tyrime dalyvavo 372 respondentai, kurie buvo atsitiktine tvarka priskirti vienam iš scenarijų. Visi tyrimo dalyviai buvo socialinių tinklų naudotojai, o tyrimas atliktas naudojant internetinės apklausos platformą „LimeSurvey“, leidusią užtikrinti atsitiktinį paskirstymą, duomenų saugumą bei etinių reikalavimų laikymąsi. Tyrimo instrumentas sudarytas remiantis anksčiau taikytomis patikimomis matavimo skalėmis, kurios buvo adaptuotos šio tyrimo kontekstui. Instrumente įtraukti pagrindiniai konstruktais grįsti kintamieji: suvokiamas šaltinio ir turinio atitikimas, suvokiamas šaltinio ir turinio patikimumas bei vartotojų įsitraukimo ketinimai, matuojami Likerto ir semantinio diferencialo skalėmis. Prieš pagrindinį tyrimą instrumentas buvo preliminariai įvertintas naudojant didelio kalbos modelio („ChatGPT“) pagalbą, siekiant patikrinti manipuliacijų loginį nuoseklumą ir kintamųjų jautrumą. Statistinei analizei buvo taikyti keli metodai: faktorinė analizė KMO ir Bartlerto testais siekiant įvertinti skalių struktūrą, Kronbacho alfa – konstrukto patikimumui, dispersinė analizė (ANOVA) ir t-testai – eksperimentinių grupių palyginimui, o koreliacijos ir sentimentų analizė – gilesnei interpretacijai. Taip pat buvo patikrinta manipuliacijos suvokimo kokybė, užtikrinant, kad respondentai tinkamai interpretavo eksperimentines sąlygas. Toks kompleksinis metodologinis pagrindimas leido užtikrinti tyrimo rezultatų validumą, patikimumą ir interpretacinį gilumą, padedantį pagrįsti hipotetizuojamus ryšius tarp DI autorystės žymėjimo, teminio turinio pobūdžio ir vartotojų reakcijų skaitmeninėje erdvėje.
4. Empiriniai tyrimo rezultatai patvirtina, kad DI turinio autorystės žymėjimas reikšmingai veikia vartotojų atsaką. Atliktas empirinis tyrimas atskleidė reikšmingas sąsajas tarp DI autorystės žymėjimo, turinio tematikos bei vartotojų suvokimo apie šaltinio ir turinio atitikimą. Rezultatai parodė, kad DI autorystės žyma turi statistiškai reikšmingą poveikį: pažymėtas turinys buvo vertinamas kaip labiau atitinkantis jį pateikusį šaltinį nei nepažymėtas. Tai dera su schemų atitikimo teorija – turinys, kurio autorystė yra aiškiai įvardyta, vartotojų sąmonėje tampa konceptualiai labiau nuoseklus bei lengviau įsisavinamas. Taip pat buvo nustatyta, kad DI sukurtas turinys, turintis sąsajų su DI, sustiprina vartotojų suvokimą apie šaltinio ir turinio atitikimą, o žymėjimo ir tematikos sąveika rodo, jog maksimalų poveikį šis ryšys įgyja tuomet,

kai turinys yra tiek susijęs su DI, tiek pažymėtas kaip DI sukurtas. Tyrimo rezultatai taip pat atskleidė nuoseklų ir prasmingą ryšį tarp vartotojų suvokiamo šaltinio ir turinio atitikimo bei šaltinio patikimumo vertinimo. Visose nagrinėtose šaltinio patikimumo dimensijose – įskaitant kompetenciją, empatiją, palankumą ir pasitikėjimą – pastebėta, kad turinys, kuris vartotojui atrodo logiškai susietas su jį pateikusių šaltiniu, skatina teigiamą šio šaltinio vertinimą. Tokį patį poveikį šaltinio ir turinio atitikimas daro ir turinio patikimumui – jei vartotojai suvokia šaltinį ir turinį kaip atitinkančius vienas kitą, turinio patikimumo vertinimas yra aukštas. Visgi, vien dermės tarp šaltinio ir turinio suvokimas nėra pakankamas veiksnys paskatinti vartotoją ištraukti į turinį socialiniuose tinkluose. Komentavimas, dalijimasis ar kitoks elgseninis atsakas tikėtinas tik tuomet, kai turinys yra įtikinamas, sąžiningas ir keliantis pasitikėjimą, o šaltinis – palankus ir empatiškas vartotojo atžvilgiu. Šie rezultatai yra svarbūs tiek akademiniam diskursui apie DI sukurtą turinio priimtinumą, tiek praktinėms komunikacijos strategijų implikacijoms, ypač socialinių tinklų aplinkoje, kur šaltinio ir turinio atitiktis tampa reikšminga auditorijos pasitikėjimo sąlyga.

Atlikus eksperimentinį tyrimą, kuriuo siekta įvertinti DI sukurtą rinkodaros turinio ir jo autorystės žymėjimo poveikį vartotojų reakcijoms socialiniuose tinkluose, parengtos praktinės rekomendacijos ir siūlomos ateities tyrimų kryptys:

- Naudoti DI technologiją kuriant turinį apie DI grįstas paslaugas ar produktus. Pažymėtas DI turinys sukelia teigiamą vartotojų atsaką tik tuomet, kai jis atitinka tematiką – pavyzdžiui, kai kalbama apie DI pagrindu sukurtas technologijas, inovacijas ar produktus. Todėl rekomenduojama DI turinio kūrimui taikyti tik tada, kai jis sustiprina komunikacinę logiką, o ne ją trikdo.
- Skatinti empatišką ir palankų bendravimą vartotojų atžvilgiu prekių ženklo komunikacijoje. Remiantis tyrimo rezultatais, vartotojų išitraukimą į DI turinį skatina suvokiamas turinio šaltinio rūpestis vartotoju, jo interesais, nesavanaudiškumas, jautrumas ir supratingumas. Todėl, net ir DI sukurtame turinyje, rekomenduojama naudoti empatišką toną, vartotojui priimtina kalbą, kad šaltinis būtų vertinamas kaip artimas ir supratingas.
- Naudojant DI technologijas turinio kūrimui, užtikrinti tekstinio turinio kalbinį tikslumą ir stilistinę natūralumą. Tyrime atlikta semantinė analizė atskleidė, jog neigiami vartotojų atsakymai dažniausiai buvo susiję su gramatinėmis, stilistinėmis ar loginėmis klaidomis tekste. Tai reiškia, kad net menkiausi kalbiniai nesklandumai gali sukelti įtarimų dėl turinio autentiškumo ar profesionalumo. Todėl DI sukurtas turinys turėtų būti peržiūrėtas ir suredaguotas prieš jį publikuojant, siekiant pašalinti netikslumus, užtikrinti logiką ir išlaikyti nuoseklų stilių, artimą natūraliai žmonių kalbai.
- Užtikrinti skaidrumą DI turinio žymėjime, ypač aukštojo mokslo institucijų komunikacijoje. Semantinės analizės rezultatai parodė, kad dalis vartotojų, susidūrę su turiniu, kuris publikuojamas aukštojo mokslo institucijos, tikisi didesnio skaidrumo ir aiškumo dėl DI naudojimo. Vienas iš respondentų pažymėjo: „Nuotrauka atrodo, lyg galėtų būt padaryta AI, bet iš KTU tikėčiausi AI label'io.“ Šis komentaras atskleidžia ne tik pasitikėjimo lygį, bet ir aukštesnius skaidrumo standartus, taikomus akademiniams ir viešosioms institucijoms. Rekomenduojama visada atvirai pažymėti DI naudojimą, taip išlaikant auditorijos pasitikėjimą ir etišką komunikaciją.
- Ateities tyrimuose būtų naudinga išplėsti tiriamų temų spektrą, įtraukiant įvairias inovacijų sritis, kad būtų galima išsamiai įvertinti DI žymėjimo bei šaltinio ir turinio atitikimo poveikį. Taip pat svarbu analizuoti turinio charakteristikas skirtingose socialinėse platformose,

siekiant geriau suprasti, kaip vartotojai sąveikauja su vizualiniais ir tekstiniais elementais ir kaip šios sąveikos skiriasi priklausomai nuo platformos specifikos (daugiau tikslių tyrimų kryptį pateikta 4.6 poskyryje).

Šios įžvalgos gali padėti organizacijoms formuoti etišką, patikimą ir į vartotojų pasitikėjimą orientuotą komunikacijos praktiką, o kartu prisidėti prie atsakingesnio DI sprendimų integravimo į viešąją skaitmeninę erdvę.

Literatūros sąrašas

1. 2024 m. birželio 13 d. Europos Parlamento ir Tarybos reglamentas (ES) 2024/1689, kuriuo nustatomos suderintos dirbtinio intelekto taisyklės ir iš dalies keičiami reglamentai (EB) Nr. 300/2008, (ES) Nr. 167/2013, (ES) Nr. 168/2013, (ES) 2018/858, (ES) 2018/1139 ir (ES) 2019/2144 ir direktyvos 2014/90/ES, (ES) 2016/797 ir (ES) 2020/1828 (Dirbtinio intelekto aktas) (2024) [žiūrėta 2024-12-30]. Prieiga per internetą: <https://eur-lex.europa.eu/legal-content/LT/TXT/?uri=CELEX:32024R1689>
2. Abi-Rafteh, J., Cattelan, L., Xu, H. H., Bassiri-Tehrani, B., Kazan, R., & Nahai, F. (2024). Artificial Intelligence-Generated Social Media Content Creation and Management Strategies for Plastic Surgeons. *Aesthetic Surgery Journal*, 44(7), 769–778.
3. Adel, M. A. A., Abouelnour, M. M., Alhourani, M. I. & Awad, A. A. (2023). Towards Intelligent Universities Enhanced with Artificial Intelligence (AI). *Journal of Infrastructure, Policy and Development* 2025, 9(1), 1–19.
4. Altay ir Gilardi, (2024). 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), 403.
5. Ananthakrishnan, R., & Arunachalam, T. (2022). Comparison Of Consumers Perception Between Human Generated And Ai Aided Brand Content. *Webology*, 19(2), 6293–6302.
6. Appelman, A., & Sundar, S. S. (2016). Measuring Message Credibility: Construction and Validation of an Exclusive Scale. *Journalism ir Mass Communication Quarterly*, 93(1), 59–79.
7. Arango, L., Singaraju, S. P., & Niininen, O. (2023). Consumer responses to AI-generated charitable giving ads. *Journal of Advertising*, 52(4), 486–503.
8. Argan, M., Dinç, H., Kaya, S., & Argan, T. M. (2023). Artificial intelligence (AI) in advertising. *Advances in Distributed Computing and Artificial Intelligence Journal*, 11(3), 331–348.
9. Baek, T. H., Kim, J., & Kim, J. H. (2024). Effect of disclosing AI-generated content on prosocial advertising evaluation. *International Journal of Advertising*, 1–22.
10. Boerman, S. C., Willemsen, L. M., & Van Der Aa, E. P. (2017). "This post is sponsored": Effects of sponsorship disclosure on persuasion knowledge and electronic word of mouth in the context of Facebook. *Journal of Interactive Marketing*, 38, 82–92.
11. Boyd, D. M., & Ellison, N. B. (2007). Social network sites: Definition, history, and scholarship. *Journal of Computer-Mediated Communication*, 13(1), 210–230.
12. Brewer, P.R., Ley, B.L. (2013). Whose Science Do You Believe? Explaining Trust in Sources of Scientific Information About the Environment. *Science Communication*, 35(1), 115–137.
13. Brodie, R. J., Hollebeek, L. D., Juric, B., & Ilić, A. (2011). Customer engagement: Conceptual domain, fundamental propositions, and implications for research. *Journal of Service Research*, 14(3), 252–271.
14. Brodie, R. J., Ilić, A., Juric, B., & Hollebeek, L. (2013). Consumer engagement in a virtual brand community: An exploratory analysis. *Journal of Business Research*, 66(1), 105–114.
15. Brondi, S., Pellegrini, G., Guran, P., Fero, M., & Rubin, A. (2021). Dimensions of trust in different forms of science communication: the role of information sources and channels used to acquire science knowledge. *Journal of Science Communication*, 20(3), 1–21.
16. Brüns, J. D., & Meißner, M. (2024). Do you create your content yourself? Using generative artificial intelligence for social media content creation diminishes perceived brand authenticity. *Journal of Retailing and Consumer Services*, 79, 103790.
17. Bui, H. T., Filimonau, V., & Sezerel, H. (2024). AI-thenticity: Exploring the effect of perceived authenticity of AI-generated visual content on tourist patronage intentions. *Journal of Destination Marketing ir Management*, 34, 100956.

18. Campbell, C., Plangger, K., Sands, S. Kietzmann, J., & Bates, K. (2022). How deepfakes and artificial intelligence could reshape the advertising industry: The coming reality of AI fakes and their potential impact on consumer behavior. *Journal of Advertising Research*, 6(23), 241–51.
19. Carney, S., Riveros, I., and Tully, S. (2024). Made with AI: Consumer Engagement with Media Containing AI Disclosures. [žiūrėta 2025-04-25]. Prieiga per internetą: <http://dx.doi.org/10.2139/ssrn.4988760>
20. Castelo, N., Bos, M. W., & Lehmann, D. R. (2019). Task-Dependent Algorithm Aversion. *Journal of Marketing Research*, 56(5), 809-825.
21. Chen, H., Wang, P., & Hao, S. (2025). AI in the spotlight: The impact of artificial intelligence disclosure on user engagement in short-form videos. *Computers in Human Behavior*, 162, 108448.
22. Cheung, C. M. K., & Thadani, D. R. (2012). The impact of electronic word-of-mouth communication: A literature analysis and integrative model. *DECISION SUPPORT SYSTEMS*, 54(1), 461–470.
23. Choi, S. (2024). Building relationships through AI-influencers: the role of AI-influencer-product congruence and interaction. *Public Relations Journal*, 17(1). [žiūrėta 2025-03-15]. Prieiga per internetą: <https://instituteforpr.org/prj-17-1-1-choi/>
24. Choi, S. M., & Rifon, N. J. (2002). Antecedents and Consequences of Web Advertising Credibility: A Study of Consumer Response to Banner Ads. *Journal of Interactive Advertising*, 3(1), 12–24.
25. Chwialkowska, A. (2019). The effects of social media content on customer engagement behavior: An exploration of informational, entertaining, and remunerative content. *Journal of Interactive Marketing*, 47, 85–96.
26. Civit, M., Civit-Masot, J., Cuadrado, F., & Escalona, M. J. (2022). A systematic review of artificial intelligence-based music generation: Scope, applications, and future trends. *Expert Systems with Applications*, 209, 118190-.
27. Clayton, K., Blair, S., Busam, J. A., Forstner, S., Glance, J., Green, G., Kawata, A., Kovvuri, A., Martin, J., Morgan, E., Sandhu, M., Sang, R., Scholz-Bright, R., Welch, A. T., Wolff, A. G., Zhou, A., & Nyhan, B. (2020). Real Solutions for Fake News? Measuring the Effectiveness of General Warnings and Fact-Check Tags in Reducing Belief in False Stories on Social Media. *Political Behavior*, 42(4), 1073–1095.
28. Cohen, J. (1988). *Statistical Power Analysis for the Behavioral Sciences* (2nd ed.). Hillsdale, NJ: Lawrence Erlbaum Associates, Publishers.
29. Cohen, J., Cohen, P., West, S. G., & Aiken, L. S. (2003). *Applied multiple regression/correlation analysis for the behavioral sciences* (3rd ed.). Lawrence Erlbaum Associates Publishers.
30. Cunningham, J. A., & Link, A. N. (2015). Fostering university-industry R&D collaborations in European Union countries. *International Entrepreneurship and Management Journal*, 11(4), 849–860.
31. Dahlén, M., Lange, F., Sjödin, H., & Törn, F. (2005) Effects of Ad-Brand Incongruence, *Journal of Current Issues & Research in Advertising*, 27(2), 1–12.
32. Davenport, T., Guha, A., Grewal, D., & Bressgott, T. (2020). How artificial intelligence will change the future of marketing. *Journal of the Academy of Marketing Science*, 48, 24–42.
33. Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340.
34. Davis, J. (2011). *Advertising research: Theory in practice*. Pearson Higher Ed.
35. De Veirman, M., & Hudders, L. (2020). Disclosing sponsored Instagram posts: the role of material connection with the brand and message-sidedness when disclosing covert advertising. *International Journal of Advertising*, 39(1), 94–130.

36. Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126.
37. Dietvorst, B. J., Simmons, J. P., & Massey, C. (2018). Overcoming algorithm aversion: people will use imperfect algorithms if they can (Even slightly) modify them. *Manage Sci*, 64(3), 1155–1170.
38. Dolan, R., Conduit, J., Fahy, J., & Goodman, S. (2016). Social media engagement behaviour: a uses and gratifications perspective. *Journal of Strategic Marketing*, 24(4), 261–277.
39. Dolan, R., Conduit, J., Frethey-Bentham, C., Fahy, J., & Goodman, S. (2019). Social media engagement behavior: A framework for engaging customers through social media content. *European Journal of Marketing*, 53(10), 2213–2243.
40. Draxler, F., Werner, A., Lehmann, F., Hoppe, M., Schmidt, A., Buschek, D., & Welsch, R. (2024). The AI Ghostwriter effect: When users do not perceive ownership of AI-generated text but self-declare as authors. *ACM Transactions on Computer-Human Interaction*, 31(2), 25.
41. Ecker, U. K. H., & Antonio, L. M. (2021). Can you believe it? An investigation into the impact of retraction source credibility on the continued influence effect. *Memory ir Cognition*, 49, 631–644.
42. Entradas, M. (2022). Public communication at research universities: Moving towards (de)centralised communication of science? *Public Understanding of Science (Bristol, England)*, 31(5), 634–647.
43. Exner, Y., Hartmann, J., Netzer, O., & Zhang, S. (2025). AI in Disguise - How AI-Generated Ads' Visual Cues Shape Consumer Perception and Performance. [žiūrëta 2025-03-15]. Prieiga per internetą: <http://dx.doi.org/10.2139/ssrn.5096969>
44. Faul, F., Erdfelder, E., Buchner, A., & Lang, A. G. (2009). Statistical power analyses using GPower 3.1: Tests for correlation and regression analyses. *Behavior Research Methods*, 41(4), 1149–1160.
45. Fox, A. K., Nakhata, C., & Deitz, G. D. (2019). Eat, drink, and create content: a multi-method exploration of visual social media marketing content. *International Journal of Advertising*, 38(3), 450–470.
46. Friestad, M., & Wright, P. (1994). The Persuasion Knowledge Model: How People Cope with Persuasion Attempts, *Journal of Consumer Research*, 21(1), 1–31.
47. Gao, B., Wang, Y., Xie, H., Hu, Y., & Hu, Y. (2023). Artificial intelligence in advertising: Advancements, challenges, and ethical considerations in targeting, personalization, content creation, and ad optimization. *Sage Open*, 13(4).
48. Geissinger, A., & Laurell, C. (2016). User engagement in social media – an explorative study of Swedish fashion brands. *Journal of Fashion Marketing and Management*, 20(2), 177–190.
49. Germelmann, C. C., Herrmann, J.-L., Kacha, M., & Darke, P. R. (2020). Congruence and Incongruence in Thematic Advertisement-Medium Combinations: Role of Awareness, Fluency, and Persuasion Knowledge. *Journal of Advertising*, 49(2), 141–164.
50. Grandinetti, J. (2023). Examining embedded apparatuses of AI in Facebook and TikTok. *AI ir Society*, 38(4), 1273–1286.
51. Grigsby, J. L., Michelsen, M., & Zamudio, C. (2025). Service ads in the era of generative AI: Disclosures, trust, and intangibility. *Journal of Retailing and Consumer Services*, 84, 104231-. [žiūrëta 2025-03-15]. Prieiga per internetą: <https://doi.org/10.1016/j.jretconser.2025.104231>
52. 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.

53. Guha, A., Grewal, D., Kopalle, P. K., Haenlein, M., Schneider, M. J., Jung, H., Moustafa, R., Hegde, D. R., & Hawkins, G. (2021). How artificial intelligence will affect the future of retailing. *Journal of Retailing*, 97(1), 28–41.
54. Hayes, A. F. (2013). Introduction to mediation, moderation, and conditional process analysis: A regression-based approach. The Guilford Press.
55. Hallock, W., Roggeveen, A. L., & Crittenden, V. (2019). Firm-level perspectives on social media engagement: An exploratory study. *Qualitative Market Research*, 22(2), 217–226.
56. Hardin, R. (2001). Conceptions and explanations of trust. In K. S. Cook (Ed.), *Trust in society*, 3–39. Russell Sage Foundation.
57. Harmon-Kizer, T. R. (2014). The effects of schema congruity on consumer response to celebrity advertising. *Journal of Marketing Communications*, 23(2), 162–175.
58. Hassan, H. E. (2024). Social Media Advertising Features that Enhance Consumers' Positive Responses to Ads. *Journal of Promotion Management*, 30(5), 874–900.
59. Haupt, M., Freidank, J., & Haas, A. (2025). Consumer responses to human-AI collaboration at organizational frontlines: strategies to escape algorithm aversion in content creation. *Review of Managerial Science*, 19(2), 377–413.
60. Hollebeek, L. D., Glynn, M. S., & Brodie, R. J. (2014). Consumer brand engagement in social media: Conceptualization, scale development and validation. *Journal of Interactive Marketing*, 28(2), 149–165.
61. Hovland, C. I., Janis, I. L., & Kelley, H. H. (1953). Communication and persuasion. Yale University Press.
62. Howard, M. C. (2023). A systematic literature review of exploratory factor analyses in management. *Journal of Business Research*, 164, 113969.
63. Huang, Y., & Yoon, H. J. (2022). Prosocial native advertising on social media: effects of ad-context congruence, ad position and ad type. *Journal of Social Marketing*, 12(2), 105–123.
64. Huang, M.-H., & Rust, R. T. (2021). A strategic framework for artificial intelligence in marketing. *Journal of the Academy of Marketing Science*, 49(1), 30–50.
65. Yang, J., Jiang, M., & Wu, L. (2021). Native Advertising in WeChat Official Accounts: How Do Ad-Content Congruence and Ad Skepticism Influence Advertising Value and Effectiveness? *Journal of Interactive Advertising*, 21(1), 17–33.
66. Yang, W. (2024). Beyond algorithms: The human touch machine-generated titles for enhancing click-through rates on social media. *PloS One*, 19(7), e0306639.
67. Yoo, K., Haenlein, M., & Hewett, K. (2024). *A Whole New World: Charting Unexplored Territories in Consumer Research with Generative Artificial Intelligence*. Working Paper, Marketing Science Institute. [žiūrėta 2025-04-29]. Prieiga per internetą: <https://www.msi.org/working-paper/a-whole-new-world-charting-unexplored-territories-in-consumer-research-with-generative-artificial-intelligence/>
68. Yoon, H. J. (2013). Understanding schema incongruity as a process in advertising: Review and future recommendations. *Journal of Marketing Communications*, 19(5), 360–376.
69. Jenkins, E. L., Ilicic, J., Barklamb, A. M., & McCaffrey, T. A. (2020). Assessing the Credibility and Authenticity of Social Media Content for Applications in Health Communication: Scoping Review. *Journal of Medical Internet Research*, 22(7), e17296–e17296.
70. Jia, Z. (2025). An experimental study on the impact of social media images on users' online social anxiety in China. *Discover Psychology*, 5(1), 14–31.
71. Johnson, V., Butterfuss, R., & Kendeou, P. (2024). Dynamic source credibility and its impacts on knowledge revision. *Memory and Cognition*, 52(7), 1548–1566.

72. Joshi, K., Loganathan, M., Chandrashekar, D., Brem, A., & Satyanarayana, K. (2024). How Do University-Based Incubators and Accelerators Build Dynamic Capabilities?-An Exploration From India. *IEEE Transactions on Engineering Management*, 71, 9035–9048.
73. Kaiser, H. F. (1974). An index of factorial simplicity. *Psychometrika*, 39(1), 31–36.
74. Kang, H., & Lou, C. (2022). *AI agency vs. human agency: Understanding human–AI interactions on TikTok and their implications for user engagement*. *Journal of Computer-Mediated Communication*, 27(5), 1–13.
75. Kaplan A. M., & Haenlein M. (2010). Users of the world, unite! The challenges and opportunities of social media. *Business Horizons*, 53(1), 59–68.
76. Kaplan, A., & Haenlein, M. (2019). Siri, Siri, in my hand: Who’s the fairest in the land? On the interpretations, illustrations, and implications of artificial intelligence. *Business Horizons*, 62(1), 15–25.
77. Katz, E., Blumler, J., & Gurevitch, M. (1973). USES AND GRATIFICATIONS RESEARCH. *Public Opinion Quarterly*, 37(4), 508–523.
78. Katz, E., Haas, H., & Gurevitch, M. (1973). On the Use of the Mass Media for Important Things. *American Sociological Review*, 38(2), 164–181.
79. Kietzmann J. H., Hermkens K., McCarthy I. P., & Silvestre B. S. (2011). Social media? Get serious! Understanding the functional building blocks of social media. *Business Horizons*, 54(3), 241–251.
80. Kim, K., & Kim, H. S. (2022). Visuals misleading consumers? Testing the visual superiority effect in advertising. *The Journal of Consumer Marketing*, 39(1), 78–92.
81. Kim, S. B., Kim, K. J., & Kim, D.Y. (2016). Exploring the effective restaurant CrM ad. *International Journal of Contemporary Hospitality Management*, 28(11), 2473–2492.
82. Kirkby, A., Baumgarth, C., & Henseler, J. (2023). To disclose or not disclose, is no longer the question: Effect of AI-disclosed brand voice on brand authenticity and attitude. *Journal of Product and Brand Management*, 32(7), 1108–1122.
83. Koch, T. K., Frischlich, L., & Lerner, E. (2023). Effects of fact-checking warning labels and social endorsement cues on climate change fake news credibility and engagement on social media. *Journal of Applied Social Psychology*, 53(6), 495–507.
84. Kohring, M., & Matthes, J. (2007). Trust in News Media: Development and Validation of a Multidimensional Scale. *Communication Research*, 34(2), 231–252.
85. Kozinets, R. V. (2014). Social brand engagement: A new idea. *Marketing Intelligence Review*, 6(2), 8–15.
86. Kudeshia, C., & Kumar, A. (2017). Social eWOM: does it affect the brand attitude and purchase intention of brands? *Management Research Review*, 40(3), 310–330.
87. Kumar, V., Rajan, B., Gupta, S., & Dalla Pozza, I. (2019). Customer engagement in service. *Journal of the Academy of Marketing Science*, 47(1), 138–160
88. Kumar, V., Rajan, B., Venkatesan, R., & Lecinski, J. (2019). Understanding the role of artificial intelligence in personalized engagement marketing. *California Management Review*, 61(4), 135–155.
89. Lange, F., & M. Dahle’n (2003). Let’s be strange: Brand familiarity and ad-brand incongruency. *Journal of Product and Brand Management*, 12(7), 449–61.
90. Laskey, H.A., Day, E. and Crask, M.R. (1989), “Typology of main message strategies for television commercials”, *Journal of Advertising*, Vol. 18 No. 1, pp. 36-41.
91. Lasswell H. D. (1948). *The structure and function of communication in society*. In Bryson L. (Ed.), *The Communication of Ideas* (37–51). New York: Harper & Row.

92. Le, T. D. (2018). Influence of WOM and content type on online engagement in consumption communities : The information flow from discussion forums to Facebook. *Online Information Review*, 42(2), 161–175.
93. Lee, D., Hosanagar, K., & Nair, H. S. (2018). Advertising Content and Consumer Engagement on Social Media: Evidence from Facebook. *Management Science*, 64(11), 5105–5131.
94. Lee, H., & Cho, C.-H. (2020). Digital advertising: Present and future prospects. *International Journal of Advertising*, 39(3), 332–341.
95. Lee, J. L., Choi, S. H., Jeong, S., & Ko, N. (2025). Generative AI in sport advertising: effects of source-message (in) congruence, model types and AI awareness. *International Journal of Sports Marketing and Sponsorship*.
96. Lee, J., & Hong, I. B. (2016). Predicting positive user responses to social media advertising: The roles of emotional appeal, informativeness, and creativity. *International Journal of Information Management*, 36(3), 360–373.
97. Lee, J., & Kim, H. (2022). How to survive in advertisement flooding: The effects of schema–product congruity and attribute relevance on advertisement attitude. *Journal of Consumer Behaviour*, 21(2), 214–230.
98. Lim, S., & Schmälzle, R. (2024). The effect of source disclosure on evaluation of AI-generated messages. *Computers in Human Behavior: Artificial Humans*, 2(1), 100058.
99. Lou, C., & Yuan, S. (2019). Influencer marketing: How message value and credibility affect consumer trust of branded content on social media. *Journal of Interactive Advertising*, 19(1), 58–73.
100. Lucassen, T., & Schraagen, J. M. (2012). Propensity to trust and the influence of source and medium cues in credibility evaluation. *Journal of Information Science*, 38(6), 566–577.
101. Luo, M., Hancock, J. T., & Markowitz, D. M. (2022). Credibility Perceptions and Detection Accuracy of Fake News Headlines on Social Media: Effects of Truth-Bias and Endorsement Cues. *Communication Research*, 49(2), 171–195.
102. Luo, X., Tong, S., Fang, Z., & Qu, Z. (2019). Machines vs. humans: The impact of artificial intelligence chatbot disclosure on customer purchases. *Journal of Marketing Research*, 56(5), 735–747.
103. MacKenzie, S. B., & Lutz, R. J. (1989). An Empirical Examination of the Structural Antecedents of Attitude toward the Ad in an Advertising Pretesting Context. *Journal of Marketing*, 53(2), 48–65.
104. Mahmud, Md. S., Islam, Md. N., Ali, Md. R., & Mehjabin, N. (2024). Impact of Electronic Word of Mouth on Customers' Buying Intention Considering Trust as a Mediator: A SEM Approach. *Global Business Review*, 25(2_suppl), S184–S198.
105. Mandler, G. (1982). The structure of value: Accounting for taste. In *Affect and cognition: The 17th annual Carnegie symposium on cognition*, ed. M.S. Clark and S.T. Fiske, 3–36. Hillsdale, NJ: Lawrence Erlbaum Associates.
106. Matz, S. C., Teeny, J. D., Vaid, S. S., Peters, H., Harari, G. M., & Cerf, M. (2024). The potential of generative AI for personalized persuasion at scale. *Scientific Reports*, 14(1), 4692–16.
107. McCay-Peet, L., & Quan-Haase, A. (2016). A model of social media engagement: User profiles, gratifications, and experiences. In *Proceedings of the 2016 Hawaii International Conference on System Sciences (HICSS)*, 3866–3875.
108. McCroskey, J. C., & Teven, J. J. (1999). Goodwill: A reexamination of the construct and its measurement. *Communications Monographs*, 66(1), 90–103.

109. Meyers-Levy, J., & Tybout, A. M. (1989). Schema Congruity as a Basis for Product Evaluation. *The Journal of Consumer Research*, 16(1), 39–54.
110. Mena, P. (2020). Cleaning Up Social Media: The Effect of Warning Labels on Likelihood of Sharing False News on Facebook. *Policy and Internet*, 12(2), 165–183.
111. Metzger, M. J., Flanagin, A. J., Eyal, K., Lemus, D. R., & Mccann, R. M. (2003). Credibility for the 21st Century: Integrating Perspectives on Source, Message, and Media Credibility in the Contemporary Media Environment. *Annals of the International Communication Association*, 27(1), 293–335.
112. Molina, M. D., Wang, J., Sundar, S. S., Le, T., & DiRusso, C. (2023). Reading, Commenting and Sharing of Fake News: How Online Bandwagons and Bots Dictate User Engagement. *Communication Research*, 50(6), 667–694.
113. Nasirov, S., & Joshi, A. M. (2023). Minding the communications gap: How can universities signal the availability and value of their scientific knowledge to commercial organizations? *Research Policy*, 52(9), 1–19.
114. Nunnally J. C. (1978). *Psychometric theory* (2nd ed.). New York, NY: McGraw-Hill.
115. Ohanian, R. (1990). Construction and Validation of a Scale to Measure Celebrity Endorsers' Perceived Expertise, Trustworthiness, and Attractiveness. *Journal of Advertising*, 19(3), 39–52.
116. Onofrei, G., Filieri, R., & Kennedy, L. (2022). Social media interactions, purchase intention, and behavioural engagement: The mediating role of source and content factors. *Journal of Business Research*, 142, 100–112.
117. Onofrei, G., Filieri, R., & Kennedy, L. (2022). Social media interactions, purchase intention, and behavioural engagement: The mediating role of source and content factors. *Journal of Business Research*, 142, 100–112.
118. OpenAI. (2023). ChatGPT [Programinė įranga] [žiūrėta 2024-11-19]. Prieiga per internetą: <https://chat.openai.com/chat>
119. Ouiridi, M. E., El Ouiridi, A., Segers, J., & Henderickx, E. (2014). Social Media Conceptualization and Taxonomy: A Lasswellian Framework. *Journal of Creative Communications*, 9(2), 107–126.
120. Pallant, J. (2020). *SPSS Survival Manual: A step by step guide to data analysis using IBM SPSS* (7th ed.). Routledge.
121. 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, 102705.
122. Pavlik, J. V. (2023). Collaborating With ChatGPT: Considering the Implications of Generative Artificial Intelligence for Journalism and Media Education. *Journalism ir Mass Communication Educator*, 78(1), 84–93.
123. Peracchio, L. A., & Tybout, A. M. (1996). The Moderating Role of Prior Knowledge in Schema-Based Product Evaluation. *The Journal of Consumer Research*, 23(3), 177–192.
124. Peres, R., Schreier, M., Schweidel, D., & Sorescu, A. (2023). On ChatGPT and beyond: How generative artificial intelligence may affect research, teaching, and practice. *International Journal of Research in Marketing*, 40(2), 269–275.
125. Piligrimienė, Ž. (2016). Marketingo tyrimų duomenų analizė SPSS programa. *Kaunas: technologija*, 7–155.
126. Pletikosa Cvijikj, I., & Michahelles, F. (2013). Online engagement factors on Facebook brand pages. *Social Network Analysis and Mining*, 3(4), 843–861.
127. Preston, C. C., & Colman, A. M. (2000). Optimal number of response categories in rating scales: reliability, validity, discriminating power, and respondent preferences. *Acta Psychologica*, 104(1), 1–15.

128. Reif, A., Taddicken, M., Guenther, L., Schroeder, J. T., & Weingart, P. (2024). The Public Trust in Science Scale: A Multilevel and Multidimensional Approach. *Science Communication*, 00(0), 1–32.
129. Rossmann, A., Ranjan, K. R., & Sugathan, P. (2016). Drivers of user engagement in eWoM communication. *Journal of Services Marketing*, 30(5), 541–553.
130. Sarstedt, M., Adler, S. J., Rau, L., ir Schmitt, B. (2024). Using large language models to generate silicon samples in consumer and marketing research: Challenges, opportunities, and guidelines. *Psychology & Marketing*, 41(6), 1254–1270.
131. Sarstedt, M., Hair, J. F., Pick, M., Liengaard, B. D., Radomir, L., ir Ringle, C. M. (2022). Progress in partial least squares structural equation modeling use in marketing research in the last decade. *Psychology & Marketing*, 39(5), 1035–1064.
132. Schidolski, C. A., Farias, D., Pickssius, M. W., De Faria Silva, R., & Da Costa Dos Santos, R. (2023). Innovative Marketing Strategies for Higher Education Institutions: A Systematic Review on the Analysis of Main Approaches and Best Practices. *IOSR Journal of Humanities and Social Science*, 28(11), 1–9.
133. Schivinski, B., Christodoulides, G., & Dabrowski, D. (2016). Measuring consumers' engagement with brand-related social-media content: Development and validation of a scale that identifies levels of social-media engagement with brands. *Journal of Advertising Research*, 56(1), 64–80.
134. Shah, S. K., & Pahnke, E. C. (2014). Parting the ivory curtain: understanding how universities support a diverse set of startups. *The Journal of Technology Transfer*, 39(5), 780–792. <https://doi.org/10.1007/s10961-014-9336-0>
135. Shamdasani, P., Stanaland, A., & Tan, J. (2001). Location, location, location: Insights for advertising placement on the web. *Journal of Advertising Research*, 41(4), 7–21.
136. Shamim, K., & Islam, T. (2022). Digital influencer marketing: How message credibility and media credibility affect trust and impulsive buying. *Journal of Global Scholars of Marketing Science*, 32(4), 601–626.
137. Shantatula, S., Lei, Z. & Wiredu, J. (2024). THE EFFECT OF ARTIFICIAL INTELLIGENCE GENERATED CONTENT AND USER GENERATED CONTENT ON SOCIAL MEDIA MARKETING STRATEGY. *International Journal of Development Research*, 14(4), 65318–65321.
138. Sitta, D., Faulkner, M., & Stern, P. (2018). What can the brand manager expect from Facebook? *Australasian Marketing Journal*, 26(1), 17–22.
139. Taylor, C. R. (2019). Editorial: Artificial intelligence, customized communications, privacy, and the General Data Protection Regulation (GDPR). *International Journal of Advertising*, 38(5), 649–650.
140. Trakimavičiūtė, G. (2017). Ryšių su klientais palaikymas pasitelkiant socialinį tinklą „Facebook“. *Informacijos Mokslai*, 77(77), 145–159.
141. Trunfio, M., & Rossi, S. (2021). Conceptualising and measuring social media engagement: A systematic literature review. *Italian Journal of Marketing*, 2021(3), 267–292.
142. van Doorn, J., Lemon, K. N., Mittal, V., Nass, S., Pick, D., Pirner, P., & Verhoef, P. C. (2010). Customer engagement behavior: Theoretical foundations and research directions. *Journal of Service Research*, 13(3), 253–266.
143. Velazquez-Solis, P. E., Ibarra-Esquer, J. E., Astorga-Vargas, M. A., Flores-Rios, B. L., Carrillo-Beltrán, M., & García Pacheco, I. A. (2024). Quantitative and Qualitative Approaches of User Engagement on Facebook Fan Page. *Trudy Instituta Sistemnogo Programirovaniâ*, 36(1), 143–156.

144. Wahid, R., Mero, J., & Ritala, P. (2023). Editorial: Written by ChatGPT, illustrated by Midjourney: generative AI for content marketing. *Asia Pacific Journal of Marketing and Logistics*, 35(8), 1813–1822.
145. Wang, P. (2019). On defining artificial intelligence. *Journal of Artificial General Intelligence*, 10(2), 1–37.
146. Wani, Z. A., & Ahmad, M. (2024). A study of social visibility and engagement of world-renowned libraries on Twitter. *Digital Library Perspectives*, 40(2), 231–243.
147. Willemsen, L. M., Neijens, P. C., & Bronner, F. (2012). The ironic effect of source identification on the perceived credibility of online product reviewers. *Journal of Computer-Mediated Communication*, 18(1), 16–31.
148. Wiredu, J., Kuma, M., Ademola, P. H., & Musesha Iyela, P. (2023). An investigation on the characteristics, abilities, constraints, and functions of artificial intelligence (AI): The age of ChatGPT as an essential ultramodern support tool. *International Journal of Development Research*, 13(5), 62614–62620.
149. Wortel, C., Vanwesenbeeck, I., & Tomas, F. (2024). Made with Artificial Intelligence The Effect of Artificial Intelligence Disclosures in Instagram Advertisements on Consumer Attitudes. *Emerging Media*, 2(3), 547–570.
150. Wu, L., & Wen, T. J. (2021). Understanding AI advertising from the consumer perspective: What factors determine consumer appreciation of AI-created advertisements? *Journal of Advertising Research*, 61(1), 58–73.
151. Wu, L., Dodoo, N. A., Wen, T. J., & Ke, L. (2021). Understanding Twitter conversations about artificial intelligence in advertising based on natural language processing. *International Journal of Advertising*, 41(4), 685–702.
152. Xie, C., Wang, Y., & Cheng, Y. (2024). Does Artificial Intelligence Satisfy You? A Meta-Analysis of User Gratification and User Satisfaction with AI-Powered Chatbots. *International Journal of Human-Computer Interaction*, 40(3), 613–623.
153. Zhang, Y., & Gosline, R. (2023). Human favoritism, not AI aversion: People’s perceptions (and bias) toward generative AI, human experts, and human-GAI collaboration in persuasive content generation. *Judgment and Decision Making*, 18.
154. Zihagh, F., Moradi, M., & Badrinarayanan, V. (2024). A brand prominence perspective on crowdfunding success for aftermarket offerings: the role of textual and visual brand elements. *The Journal of Product and Brand Management*, 33(1), 91–107.

Informacijos šaltinių sąrašas

1. Ceci, L. (2025). *Distribution of YouTube users worldwide as of February 2025, by age group and gender*. [žiūrėta 2025-05-03]. Prieiga per internetą: <https://www.statista.com/statistics/1287137/youtube-global-users-age-gender-distribution/>
1. Clegg, N. (2024). Labeling AI-Generated Images on Facebook, Instagram and Threads. Meta. [žiūrėta 2024-11-17]. Prieiga per internetą: <https://about.fb.com/news/2024/02/labeling-ai-generated-images-on-facebook-instagram-and-threads/>
2. Dixon J. S. (2024a). Global marketers on AI generated social media content labels 2024, Statista [žiūrėta 2025-01-12]. Prieiga per internetą: <https://www-statista-com.ezproxy.ktu.edu/statistics/1489699/global-marketers-ai-generated-labels-social-media/>
3. Dixon J. S. (2024b). Reactions of adults in the United Kingdom (UK) to seeing content labelled as AI generated on social media as of May 2024, Statista [žiūrėta 2025-03-15]. Prieiga per internetą: <https://www.statista.com/statistics/1489707/uk-reactions-ai-generated-content-on-social-media/>
4. Dixon J. S. (2024c). *Distribution of Facebook users worldwide as of April 2024, by age and gender*. [žiūrėta 2025-05-03]. Prieiga per internetą: <https://www.statista.com/statistics/376128/facebook-global-user-age-distribution/>
5. Dixon J. S. (2024d). *Distribution of Instagram users worldwide as of April 2024, by age and gender*. [žiūrėta 2025-05-03]. Prieiga per internetą: <https://www.statista.com/statistics/248769/age-distribution-of-worldwide-instagram-users/>
6. Fried, I. (2024). OpenAI says ChatGPT usage has doubled since last year. Axios. [žiūrėta 2024-11-17]. Prieiga per internetą: <https://www.axios.com>
7. Oficialios statistikos portalas (2021). [žiūrėta 2025-04-05]. *Lietuvos Respublikos 2021 m. gyventojų ir būstų surašymo rezultatai*. Prieiga per internetą: <https://osp.stat.gov.lt/2021-gyventoju-ir-bustu-surasymo-rezultatai/issilavinimas>
8. OpenAI (n.d.). DALL·E 2. [žiūrėta 2025-02-09]. Prieiga per internetą: <https://openai.com/index/dall-e-2/>
9. Statista Research Department (2024). *Most liked features of social media ads USA 2024: Features consumers from the United States liked most about advertising on entertainment apps or social media platforms as of May 2024*. [žiūrėta 2025-02-13]. Prieiga per internetą: <https://www-statista-com.ezproxy.ktu.edu/statistics/1489632/features-ads-usa/>
10. Survey monkey Inc. (2024). Marketing trends for a new era. Report. [žiūrėta 2024-11-24]. Prieiga per internetą: <https://www.surveymonkey.com/resources/premium/marketing-trends-for-a-new-era/>
11. TikTok (n.d.). About AI-generated content. TikTok Help Center. [žiūrėta 2024-11-17]. Prieiga per internetą: <https://support.tiktok.com/en/using-tiktok/creating-videos/ai-generated-content>

Priedai

1 priedas. ChatGPT sugeneruotas tekstas

Pateikus užklausą „Trumpai (100 ž) paaiškink man, kaip veikia „ChatGPT“? Kokių būdu, pasitelkiant tavo pagalbą, marketingo specialistai gali kurti rinkodaros turinį?“, ChatGPT sugeneravo atsakymą:

ChatGPT veikia naudodamas dirbtinio intelekto modelį, išmokytą analizuoti didelius teksto duomenų rinkinius. Modelis generuoja atsakymus remdamasis tikimybiniais ryšiais tarp žodžių ir frazių, taip imituodamas žmogaus kalbą. Tai leidžia kurti tekstus, atsakyti į klausimus ir net generuoti kūrybinį turinį.

Marketingo specialistai gali naudoti ChatGPT kurdami turinį įvairiais tikslais: generuoti idėjas kampanijoms, rašyti įrašus socialiniams tinklams, kurti SEO draugiškus tinklaraščio įrašus, laiškus ar reklamų tekstus. ChatGPT padeda greitai parengti skirtingus *tekstų variantus*, optimizuoti turinį ir taupyti laiką. Be to, modelis leidžia personalizuoti pranešimus pagal tikslią auditoriją ir rinkos tendencijas.

2 priedas. Išsami mokslinių šaltinių, aprašančių vartotojų požiūrį į DI valdomų sprendimų sukurtą rinkodaros turinį, apžvalga.

Autorius (-iai), metai	Tyrimo objektas	Teorijos	Tyrimo tikslai / hipotezės	Metodai	Tyrimo rezultatai	Siūlomos tolimesnių tyrimų kryptys
Wu et al., 2021	Įrašai, „Twitter“ platformoje apie DI naudojimą rinkodaroje	Reklamos procesų teorija, kuri apima tyrimuose nagrinėjamas reklamos sritis, įskaitant auditorijas, priemones, žiniasklaidos kanalus, pranešimus, pranešimų šaltinius, kontekstus ir reklamos organizacijas.	Išsiaiškinti, kokiomis temomis diskutuojama apie DI reklamą „Twitter“ bei kokios yra šių diskusijų nuotaikos	Natūralios kalbos apdorojimas, temų modeliavimas, sentimentų analizė	Pagrindinės temos: DI reklamos tikslinis orientavimas, žmogaus ir DI sąveika, DI įrankiai, DI turinio kūrimas. DI įrankiai, skirti reklamos efektyvumui didinti, buvo vertinami teigiamai. Tačiau DI įtraukimas į socialinės medijos kampanijas buvo vertinamas neigiamai dėl galimų privatumo pažeidimų, manipuliacijos ir autentiškumo stokos.	Siūloma tyrinėti kitas socialines platformas ir visuomenės reakciją į DI reklamą skirtinguose kultūriniuose kontekstuose.
Wu & Wen, 2021	Veiksniai, lemiantys vartotojų požiūrį į DI sukurtas reklamas	Tyrimas rėmėsi žmogaus ir DI sąveikos teorija bei „machine heuristic“ ir „Uncanny valley“ modeliais.	Nustatyti, kaip reklamos kūrimo objektyvumas, DI žmogiškumo suvokimas ir vartotojų jaučiamas nejaukumas robotų atžvilgiu daro įtaką vartotojų požiūriui į DI sukurtą rinkodaros turinį. Hipotezės: (H ₁) „Mašininė euristika“ (DI kaip patikimo ir objektyvaus suvokimas) teigiamai veikia vartotojų požiūrį. (H ₂) Suvokiamas DI reklamos „keistumas“ (vartotojų diskomfortas dėl nenatūralumo) neigiamai	Internetinė apklausa su 528 vartotojais iš Jungtinių Amerikos Valstijų, naudojant struktūrinių lygčių modelį (SEM) analizei. Tyrime įvertintos įvairios dimensijos, tokios kaip „mašininė euristika“, suvokiamas keistumas ir reklamos objektyvumas.	Tyrimo rezultatai parodė, kad tais atvejais, kai vartotojai mano, kad reklama sukurta objektyviai, jie ją vertina palankiau. Tačiau jei žmonės jaučia nejaukumą dėl robotų ar DI technologijų, tai gali turėti dvilypį poveikį – iš vienos pusės, sustiprinti pasitikėjimą DI (kaip tikslu ir patikimu įrankiu), bet tuo pačiu ir sukelti nerimą, kuris mažina reklamos patrauklumą.	Autoriai rekomendavo toliau tirti, kaip DI sukurtų reklamų dizainas gali mažinti vartotojų diskomfortą, pavyzdžiui, pabrėžiant reklamos kūrimo objektyvumą ar mažinant žmogiškumą imituojančius elementus. Jie taip pat siūlė gilintis į kultūrinius skirtumus ir jų įtaką vartotojų reakcijoms į DI technologijas, ypatingą dėmesį skiriant skirtingų amžiaus grupių bei socialinių sluoksnių požiūriui.

			veikia jų požiūrį. (H ₃) Suvokiamas reklamos kūrimo objektyvumas teigiamai veikia euristiką, bet mažina keistumą.			
Argan et al., 2022	Vartotojų elgesys reaguojant į DI sukurtas reklamas socialinėje medijoje.	Tyrimas rėmėsi Stimulo-Organizmo-Reakcijos (SOR) teorija.	Nustatyti, kaip socialinės medijos vartotojai reaguoja į DI sukurtas reklamas. Hipotezės: (H ₁) Vartotojai teigiamai reaguoja į DI reklamas, kurios sukuria susidomėjimą, pavyzdžiui, pasitelkiant elementus, skatinančius baimę praleisti (angl. fear of missing out, FoMO), asmeninius interesus ar informacijos gavimą; (H ₂) DI reklamos, kurios taupo laiką ir pateikia naujas idėjas sulaukia daugiau vartotojų dėmesio; (H ₃) Reklamos, kurios neatitinka vartotojų lūkesčių arba kelia neaiškumų, yra vertinamos neigiamai.	Kokybinis tyrimas su 23 socialinės medijos vartotojais naudojant pusiau struktūruotus interviu.	Tyrimas atskleidė tris vartotojų reakcijų į DI reklamas etapus: priėmimas (teigiamos ir neigiamos reakcijos), išitraukimas (galimybė palyginti produktus / paslaugas neieškant papildomos informacijos, laiko taupymas, papildomų žinių suteikimas) ir galutinė nuomonė (pozityvi ar negatyvi). Reklamos, kurios atitinka vartotojų lūkesčius ir suteikia naujos, vartotojams nežinomos, bet naudingos informacijos, yra geriau vertinamos, o neatitinkančios lūkesčių ar per daug abstrakčios – atmetamos.	Ateities tyrimuose buvo siūloma įtraukti didesnes ir kultūriškai įvairesnes imtis, taip pat nagrinėti, kaip DI gali gerinti vartotojų patirtį dinamiškose situacijose.
Kirkby et al., 2023	Prekių ženklo „Adidas“ rinkodaros tekstai, kurie atskleidžiami kaip sukurti DI, sukurti žmogaus arba jų autorystė nėra atskleidžiama.	Tyrimas remiasi algoritmų vengimo teorija, kuri teigia, kad vartotojai skeptiškai žiūri į algoritmų sukurtus sprendimus, net jei jie yra tikslūs. Be to, buvo panaudota prekių ženklo autentiškumo teorija, kuri akcentuoja tęstinumą, natūralumą ir patikimumą	Įvertinti, ar DI sukurto turinio autorystės atskleidimas daro poveikį suvokiamam prekių ženklo balso autentiškumui, bendram turinio autentiškumui ir vartotojų požiūriui į jį. Hipotezės: (H _{1a}) DI autorystės atskleidimas mažina prekių ženklo balso autentiškumą;	Eksperimentinis dizainas ir apklausa: (1) Teksto emocingumo lygis (žemas, vidutinis, aukštas); (2) Teksto autorystės atskleidimas (DI, žmogus, nežinoma). Dalyviai (N = 624) skaitė tekstus apie „Adidas“	Tyrimas parodė, kad DI sukurtas turinys nėra suvokiamas kaip mažiau autentiškas nei žmogaus sukurtas turinys. Požiūris į prekių ženklą taip pat reikšmingai nesikeičia net jei turinio autorystė yra atskleidžiama. Teksto emocingumo lygis (žemas, vidutinis, aukštas) neturėjo	Autoriai rekomendavo toliau tirti, kaip DI ir žmogaus bendradarbiavimas galėtų paveikti prekių ženklo kalbos autentiškumą ir plėsti dalyvių demografiją, įtraukiant kitų kultūrų ir amžiaus grupių atstovus. Taip pat siūloma gilintis į interaktyvius dvikrypčius

		kaip autentiškumo pagrindinius elementus.	(H _{1b}) DI autorystės atskleidimas neturi įtakos prekių ženklo balso autentiškumui; (H _{2a}) DI autorystės atskleidimas mažina teigiamą požiūrį į prekių ženklą; (H _{2b}) DI šaltinio atskleidimas neturi įtakos požiūriui į prekių ženklą.	produktus, o tekstų autorystė buvo nenurodyta arba nurodyta kaip priklausanti DI ar žmogui.	reikšmingos įtakos vartotojų požiūriui.	vartotojo ir DI komunikacijos procesus.
Zhang & Gosline, 2023	Žmogaus, DI ar žmogaus ir DI sukurtas turinys reklamose	Algoritmų vengimo, žmogaus favoritizmo teorijos	Tyrimo tikslai: 1. Ištirti, kaip žmonės vertina turinį, sukurtą žmogaus, DI, ar žmogaus ir DI bendradarbiaujant. 2. Nustatyti, ar žmonės rodo favoritizmą žmogui ar vengia DI. Hipotezės: H1: Žmonės geriau vertins turinį, pažymėtą kaip sukurtą žmogaus. H2: DI sukurtas turinys, jei nepažymėtas, bus įvertintas teigiamai. H3: Žmogaus favoritizmas lemia šį šališkumą, o ne neigiamas požiūris į DI.	Penki eksperimentai su dalyviais (N = 1203), tyrimas skirtingomis sąlygomis: 1. Tik turinio kokybė, be šaltinio atskleidimo. 2. Dalinai informuoti apie šaltinį. 3. Pilnai informuoti apie šaltinį.	DI kurtas turinys buvo vertinamas teigiamai, kai autorystė nebuvo nurodyta. Dalyviai teikė pirmenybę turiniui, pažymėtam kaip žmogaus kurtam, net jei DI turinys buvo objektyviai geresnis. Nepažymėtas DI turinys gavo aukštesnius įvertinimus nei pažymėtas. Favoritizmas žmogaus naudai buvo labiau susijęs su socialinėmis normomis nei su DI vengimu.	Autoriai siūlo tirti kultūrinius skirtumus ir jų poveikį DI turinio vertinimui, toliau analizuoti, kaip sumažinti favoritizmą žmogaus naudai.
Shantatula et al., 2024	DI sugeneruotas turinys (AIGC) ir vartotojų sugeneruotas turinys (UGC) socialinės medijos rinkodaroje	Technologijų priėmimo modelių (TAM) (Davis, 1989). Ši teorija aiškina, kaip vartotojai priima ir naudoja technologijas, vertindami jų naudingumą bei naudojimo paprastumą. Straipsnyje TAM modelis pritaikytas aiškinant, kaip	Ištirti, kaip DI sugeneruotas turinys ir vartotojų sugeneruotas turinys veikia vartotojus, identifikuoti jų privalumus, trūkumus. Hipotezės: (H ₁) DI generuotas turinys turi reikšmingą poveikį socialinės medijos rinkodaros strategijai. (H ₂)	Kiekybinis tyrimas: anketinė apklausa (N = 111) ir kokybiniai interviu su ekspertais.	DI naudojimas leidžia personalizuoti ir greitai sukurti didelį kiekį turinio, tačiau, vartotojų požiūriu, toks turinys stokoja autentiškumo ir kūrybiškumo. Vartotojų sugeneruotas turinys yra suprantamas kaip patikimesnis ir	Autoriai rekomenduoja toliau nagrinėti, kaip efektyviai integruoti AIGC ir UGC, siekiant užtikrinti optimalią turinio kokybę ir prekių ženklo autentiškumą. Be to, būtina gilintis į tai, kaip vartotojų kultūriniai skirtumai ar socialinių tinklų algoritmų

		DI sukurtu turinio personalizavimas ir naudotojų patirtis skatina vartotojų įsitraukimą. Tuo tarpu UGC efektyvumą TAM aiškina per vartotojų polinkį pasitikėti kitų žmonių, o ne prekių ženklų, kuriu turiniu.	Vartotojų sugeneruotas turi reikšmingą poveikį socialinės medijos rinkodaros strategijai.		įtikinamesnis. Autoriai rekomenduoja įsivertinti abiejų strategijų naudojimo privalumus ir trūkumus, mat tai gali stiprinti vartotojų įsitraukimą ir lojalumą.	pokyčiai veikia AIGC ir UGC sąveiką.
Lim & Schmalzle, 2024	Autorystės atskleidimo poveikis DI sukurtų ir žmogaus sukurtų sveikatos komunikacijos pranešimų vertinimui	Tyrime remtasi įtaigos teorija, kuri pabrėžia autorystės svarbą įtikinimo procese, bei nagrinėjančia, kaip autoriaus savybės (pavyzdžiui, ar pranešimą sukūrė DI, ar žmogus) keičia auditorijos suvokimą ir reakcijas. Be to, buvo nagrinėjamos bendros nuostatos apie DI, kurios galėjo moderuoti autorystės poveikį.	Tyrimo tikslai: 1. Nustatyti, kaip autorystės atskleidimas veikia pranešimų vertinimą. 2. Ištirti, kaip žmonių nuostatos dėl DI moderuoja šį efektą. Hipotezės: H1: Dalyviai žemesniais balais vertins DI kurtus pranešimus, kai bus nurodyta, kad juos kūrė DI. H2: Dalyviai pirmenybę teiks žmogaus kurtam turiniui, jei žinos autorių. H3: Neigiama nuostata dėl DI moderuos autorystės atskleidimo įtaką pranešimų vertinimui. H4: Neigiama nuostata dėl DI moderuos autorystės atskleidimo poveikį pasirinkimui.	Du eksperimentai: pirmame dalyvavo 142 jauni suaugę (18-24 metų), antrame – 183 įvairaus amžiaus suaugę; naudojami likerto skalės vertinimai, regresijos modeliai.	Pirmame tyrime nustatyta, kad dalyviai vertino DI sukurtus pranešimus geriau, kai autorystė nebuvo atskleista. Atskleidus autorystę, reikšmingo skirtumo tarp žmogaus ir DI kurtų pranešimų nebuvo. Antrojo tyrimo metu neigiama nuostata dėl DI moderavo autorystės poveikį: neigiamas požiūris į DI buvo susijęs su didesniu DI kurtų pranešimų efektyvumo vertinimu.	Siūloma tirti DI ir žmogaus sukurtų pranešimų vertinimo kontekstus skirtingose kultūrose ir srityse. Taip pat gilintis į autorystę nurodančias etiketes ir jų poveikį auditorijos pasitikėjimui.

Gu et al., 2024	DI įrankių sukurtų reklamų savybių (tikroviškumo, gyvybingumo, kūrybiškumo & sintetiškumo) poveikis vartotojų suvokimui ir priimtinumui	Tyrimas rėmėsi Stimulo-Organizmo-Reakcijos (SOR) teorija, kuri pabrėžia, kad įvairūs išoriniai stimulai lemia vartotojų psichologinę būseną, o ši lemia elgesio reakcijas.	Nustatyti, kaip DI sukurtų reklamų savybės (tikroviškumas, gyvybingumas, kūrybiškumas & sintetiškumas) veikia vartotojų suvokimą apie reklamos keistumą, intelektą ir priimtinumą. Hipotezės: (H ₁) Tikroviškumas mažina suvokiamą keistumą ir didina intelektą; (H ₂) Gyvybingumas mažina suvokiamą keistumą ir didina intelektą; (H ₃) Kūrybiškumas mažina suvokiamą keistumą ir didina intelektą; (H ₄) Sintetiškumas didina suvokiamą keistumą ir mažina intelektą; (H ₅) Suvokiamas keistumas mažina norą priimti reklamas; (H ₆) Suvokiamas intelektas didina reklamos priimtinumą.	Internetinė apklausa su 1147 Kinijos vartotojais.	Tikroviška, suvokiama kaip gyvybinga / energinga ir kūrybinga reklama didina vartotojų pasitikėjimą DI ir skatina palankumą matomai reklamai. Sintetiškumo jausmas sukelia nejaukumą, kuris mažina vartotojų norą priimti reklamą.	Tyrimo autoriai siūlė toliau tirti, kaip kultūriniai skirtumai gali paveikti DI sukurtų reklamų suvokimą, nes šis tyrimas buvo vykdytas tik Rytų kultūroje. Be to, buvo rekomenduojama išplėsti modelį, įtraukiant kitus išorinius veiksnius, galinčius paveikti vartotojų reakcijas.
Park et al., 2024	„Instagram“ naudotojų požiūris į DI, nuomonės formuotojų ir paprastų vartotojų paskyras bei jose skelbiamą turinį	Tyrimo metu nenurodyta	Nustatyti, kaip vartotojai suvokia DI sukurtą turinį, ar atpažįsta jį. Tyrimo klausimai: 1. Ar naudotojų prioritetai skiriasi pagal paskyros tipą? 2. Kaip skiriasi turinio kokybės vertinimai? 3. Ar naudotojai gali atpažinti DI turinį?	Anketinės apklausos, kiekybinė analizė	Respondentai dažnai negalėjo atskirti DI sukurtą turinį nuo žmogaus kuriamo, tačiau nuostatos apie turinio kokybę ir patikimumą skyrėsi priklausomai nuo paskyros tipo. Nuomonės formuotojų paskyros buvo vertinamos kaip patikimiausios ir profesionaliausios, tačiau	Siūloma plėtoti tyrimus, nagrinėjant DI įtaką kitoms socialinių tinklų platformoms ir skirtingų kultūrinių bei demografinių grupių reakcijas. Taip pat autoriai rekomenduojama gilintis į etinius DI naudojimo aspektus ir ugdyti visuomenės

					mažiau autentiškos nei paprastų vartotojų. DI paskyros buvo vertinamos kaip inovatyvios, bet jomis stokota pasitikėjimo.	gebėjimą kritiškai vertinti DI sukurtą turinį.
Lee et al., 2025	DI sukurtų sporto reklamų vertinimas: kaip reklamos modelio tipas, DI autorystės atskleidimo momentas ir šaltinio bei žinutės (ne)atitiktis veikia vartotojų požiūrį į reklamą	Šaltinio įtakos teorija (angl. source effect): reklamos šaltinio (pavyzdžiui, DI ar žmogaus) savybės (patikimumas, panašumas, patrauklumas) lemia įtikinamumą. Šaltinio ir turinio (ne)atitikties teorija: neatitikimas tarp žinutės pobūdžio (pavyzdžiui, emocinės sporto temos) ir šaltinio (pavyzdžiui, racionalaus DI) mažina reklamos efektyvumą. „Uncanny Valley“ efektas: kai DI modeliai atrodo per daug panašūs į žmones, bet vis tiek dirbtini – tai sukelia diskomfortą.	Tyrimas siekė išsiaiškinti, kaip skirtingas DI autorystės atskleidimo momentas (prieš / po / nežinoma), reklamos modelio tipas (žmogus, virtualus modelis, skaitmeninis dvynys) ir šaltinio–žinutės atitiktis paveikia vartotojų reklamos vertinimą. Hipotezės: (H ₁) Reklamos vertinimas priklauso nuo DI autorystės atskleidimo laiko; (H ₂) Vertinimas priklauso nuo modelio tipo (žmogus, virtualus, skaitmeninis dvynys); (H ₃) Šaltinio ir žinutės neatitiktis neigiamai veikia reklamos vertinimą; (H ₄) DI autorystės atskleidimo momentas ir modelio tipas interaktyviai veikia vertinimą esant neatitiktčiai.	Tyrimo taikytas 2 × 3 × 3 tarpgrupinis eksperimentinis dizainas, kuriame dalyvavo 229 Pietų Korėjos universitetų studentai. Dalyviai žiūrėjo Adidas prekių ženklų sporto reklamas, kuriose modelis demonstruodavo sportinę veiklą, o autorystės žymėjimas buvo pateikiamas skirtingu metu – prieš arba po peržiūros.	Tyrimas parodė, kad reklamos, kuriose DI autorystė buvo atskleidžiama po reklamos peržiūros, buvo vertinamos palankiausiai, o vertinimai buvo ženkliai žemesni, kai autorystė buvo atskleidžiama prieš peržiūrą. Virtualūs modeliai, kai žiūrovai nežinojo, kad jie sukurti DI, buvo vertinami prasčiausiai, tačiau, kai DI autorystė buvo atskleista po reklamos, šie modeliai buvo vertinami geriau nei tradiciniai žmonių ar skaitmeninių dvynių modeliai.	Autoriai siūlo ateityje nagrinėti, kaip kultūriniai skirtumai gali paveikti DI reklamos vertinimą, kadangi šiame tyrime dalyvavo tik Pietų Korėjos studentai. Taip pat rekomenduojama tirti, kaip pakartotinis reklamos stebėjimas veikia vartotojų požiūrį, nes realiame gyvenime vartotojai reklamą mato daugiau nei vieną kartą.

priedas. Išsami mokslinių šaltinių, aprašančių vartotojų atsaką į DI valdomų sprendimų sukurtą rinkodaros turinį, apžvalga.

Autorius (-iai), metai	Tyrimo objektas	Teorijos	Tyrimo tikslai / hipotezės	Metodai	Tyrimo rezultatai	Siūlomos tolimesnių tyrimų kryptys
------------------------	-----------------	----------	----------------------------	---------	-------------------	------------------------------------

Luo et al., 2019	DI pokalbių robotų identiteto atskleidimo poveikis klientų pirkimo sprendimams	Tyrimas susijęs su antropomorfizmo teorijomis, tačiau nenurodo konkrečios teorijos, kuria remiasi.	Ištirti, kaip DI pokalbių robotų identiteto atskleidimas (anksčiau ar vėliau pokalbyje) veikia klientų pirkimo sprendimus ir suvokimą apie pokalbių robotus. Hipotezė: DI pokalbių robotų efektyvumas priklauso nuo atskleidimo strategijos, o vėlyvas atskleidimas gali sumažinti neigiamą poveikį.	Eksperimentas realioje aplinkoje buvo atliktas finansinių paslaugų įmonėje, siekiant stebėti klientų elgesį reaguojant į DI pokalbių robotus. Imtis: 6255	Tyrimas naudojo eksperimentinį dizainą su realiais pirkimų duomenimis ir parodė, kad neatskleidžiant pokalbių robotų tapatybės pasiekiamas didesnis jų efektyvumas, tačiau jų identiteto atskleidimas dar prieš vartotojo ir pokalbių robotų sąveiką sumažina pirkimus 79,7 proc.. Kuo vėliau tapatybė yra atskleidžiama, tuo neigiamas poveikis yra mažesnis.	Autoriai siūlo tirti pokalbio tarp pokalbių roboto ir kliento dinaminius skirtumus lyginant su pokalbiu tarp darbuotojo ir kliento. Taip pat, tikrinti rezultatų tinkamumą kitose situacijose.
Powers et al., 2023	Išmaniųjų vaizdo klastočių autorystės atskleidimo poveikis vartotojų pirkimo ketinimams	Schemų suderinamumo teorija, įtaigos žinių modelis	Ištirti, kaip išmaniųjų vaizdo klastočių naudojimo atskleidimas veikia vartotojų požiūrį į reklamos šaltinio patikimumą ir jų pirkimo ketinimus. H1: Išmaniųjų vaizdo klastočių autorystės atskleidimas sumažins vartotojų pirkimo ketinimus, palyginti su neatskleistu turiniu. H2a: Išmaniųjų vaizdo klastočių autorystės atskleidimas sumažins šaltinio patikimumo suvokimą. H2b: Išmaniųjų vaizdo klastočių autorystės atskleidimas sumažins šaltinio fizinio patrauklumo suvokimą. H2c: Išmaniųjų vaizdo klastočių autorystės	Eksperimentinis dizainas su 318 dalyvių iš JAV ir Indijos: Dalyviai buvo atsitiktinai suskirstyti į dvi grupes – su atskleidimu ir be jo. Stimuliuojami vaizdo įrašai apie mobiliųjų programėlių prekybos akcijomis paslaugas.	Išmaniųjų vaizdo klastočių atskleidimas sumažino vartotojų pirkimo ketinimus ir šaltinio patikimumo suvokimą, tačiau fizinio patrauklumo ir ekspertiškumo suvokimas reikšmingai nesikeitė. Šaltinio patikimumas reikšmingai darė įtaką pirkimo ketinimams, ypač moterų grupėje. Socialinis cinizmas neturėjo reikšmingo poveikio. H5 hipotezių rezultatai parodė, kad JAV vartotojai buvo jautresni šaltinio patikimumo ir ekspertiškumo pokyčiams nei Indijos vartotojai.	Siūloma giliau analizuoti socialinius ir kultūrinius skirtumus, turinčius įtakos išmaniųjų vaizdo klastočių naudojimui rinkodaroje. Taip pat rekomenduojama tirti, kaip vartotojai suvokia sudėtingesnę ar mažiau įtikinamą informaciją, kai informacijos šaltinis yra DI.

			atskleidimas sumažins šaltinio ekspertiskumo suvokimą. H3: Šaltinio patikimumas veiks ryšį tarp atskleidimo ir pirkimo ketinimų. H4: Socialinis cinizmas sustiprins atskleidimo poveikį šaltinio patikimumui ir pirkimo ketinimams. H5a: Išmaniųjų vaizdo klastočių atskleidimo poveikis šaltinio patikimumui bus stipresnis tarp JAV vartotojų nei Indijos vartotojų. H5b: Išmaniųjų vaizdo klastočių atskleidimo poveikis šaltinio fiziniam patrauklumui bus stipresnis tarp JAV vartotojų nei Indijos vartotojų. H5c: Išmaniųjų vaizdo klastočių atskleidimo poveikis šaltinio ekspertizei bus stipresnis tarp JAV vartotojų nei Indijos vartotojų.			
Arango et al., 2023	Vartotojų reakcijos į DI sugeneruotus reklaminius vaizdus labdaros komunikacijoje	Straipsnyje remiamasi vertybių teorijos pateikta motyvų klasifikacija, kuri išskiria vidinius ir išorinius bei etinius ir instrumentinius motyvus.	Ištirti, kaip vartotojai reaguoja į DI sugeneruotus vaizdus labdaros reklamoje, ypatingą dėmesį skiriant emocijoms, empatijai, kaltės jausmui ir aukojimo ketinimams. Hipotezės: H1: Netikrumo suvokimas (DI generuotų vaizdų) mažina empatiją. H2: Empatija teigiamai	Trijų eksperimentų analizė: 1. DI generuotų ir realių vaizdų poveikio vartotojų emocijoms analizė. 2. DI vaizdų motyvų (etinių & instrumentinių) poveikio tyrimas. 3. DI vaizdų	Vaizdų netikrumo suvokimas mažina empatiją, kaltės jausmą ir aukojimo ketinimus. Tyrimo dalyviai geriau priėmė DI generuotus vaizdus, kai buvo pabrėžiama, kad jie naudojami dėl etinių priežasčių (pavyzdžiui, vaikų privatumo apsaugos).	Autoriai rekomenduoja toliau tirti skirtingoms demografinėms grupėms priklausančius, tam tikrus asmenybės bruožus, įsitikinimus turinčius asmenis mat jie galėtų skirtingai reaguoti į tokias reklamas.

			veikia kaltės jausmą. H3: Kaltės jausmas teigiamai veikia aukojimo ketinimus. H4: Netikrumo suvokimo neigiamą poveikį aukojimo ketinimams mažina empatija ir kaltės jausmas. H5: Empatija teigiamai veikia emocijų suvokimą. H6: Emocijų suvokimas teigiamai veikia aukojimo ketinimus. H7: Netikrumo suvokimo poveikį aukojimo ketinimams mažina empatija ir emocijų suvokimas.	priimtino ekstremaliomis situacijomis analizė (pavyzdžiui, nelaimių metu).	Instrumentinių motyvų akcentavimas (pavyzdžiui, taupymo) buvo mažiau veiksmingas arba net pakenkdavo aukojimo ketinimams. Nelaimės atvejais vartotojai mažiau reaguodavo į netikrumo suvokimą. DI vaizdai buvo priimami beveik taip pat kaip tikri vaizdai, nes suvokta, kad kitos alternatyvos buvo neprieinamos.	Taip pat, ilgalaikiai tyrimai siekiant nustatyti, kaip visuomenės požiūris keičiasi visuomenės požiūris į DI reklamas, kai jos vis plačiau naudojamos, galėtų suteikti vertingų įžvalgų.
Baek et al., 2024	DI turinio autorystės atskleidimo poveikis labdaros reklamoje	Įtaigos žinių modelis (angl. persuasion knowledge model), nurodantis, kad kuo asmuo daugiau žino apie įvairias įtikinimo taktikas, tuo jis yra mažiau paveiklus įtaigoms žinutėms, todėl geba atsisipirti reklamoms. Straipsnyje minimos ir algoritmo atmetimo bei antropomorfizmo teorijos.	H1: Reklamos, kuriose atskleidžiama, kad turinį sukūrė DI, sulaukia prastesnio vertinimo ir mažesnių aukojimo ketinimų nei reklamos be tokios informacijos. H2: Reklamos patikimumo suvokimas daro įtaką DI atskleidimo poveikiui ir požiūriui į reklamą ir aukojimo ketinimams. H3: DI suvokimas, atsižvelgiant į jo panašumą į žmogų, švelnina neigiamą DI atskleidimo poveikį, o jo suvokimas kaip mažiau panašaus į žmogų, jį stiprina.	Eksperimentiniai tyrimai, kiekybinė analizė	Reklamos, kuriose atskleista, kad turinį sukūrė DI, buvo vertinamos kaip mažiau patikimos, jų bendras vertinimas buvo prastesnis, o aukojimo ketinimai mažesni. Dalyviai, kurie DI suvokė kaip labiau „žmogišką“, reagavo mažiau neigiamai, o tie, kurie jį suvokė kaip „mašininį“, turėjo itin neigiamą požiūrį ir mažiausią norą aukoti.	Siūloma tyrinėti DI turinio atskleidimo poveikį kitose kultūrinėse ir demografinėse grupėse, įvairiuose medijų formatuose bei komercinėje reklamoje.
Bui et al., 2024	DI sukurtų vizualizacijų suvokiamas	Suplanuotos elgsenos teorija (TPB), autentiškumo teorijos	Tyrimo tikslai: 1. Ištirti, kaip DI sukurtų vizualizacijų suvokiamas	Eksperimentiniai tyrimai (N = 200, N = 400); regresijos	DI vizualizacijų autentiškumas didina pasitikėjimą ir ketinimus	Siūloma tirti skirtingus kultūrinius kontekstus, DI turinio poveikį kitoms

	autentiškumas turizmo srityje		autentiškumas veikia turistų ketinimus lankyti tam tikrose vietose. 2. Įvertinti pasitikėjimo ir suderinamumo vaidmenį tarp autentiškumo ir turistų ketinimų. Hipotezės: H1: DI sukurtų vizualizacijų suvokiamas autentiškumas didina turistų ketinimus lankyti. H2: DI vizualizacijų autentiškumas didina pasitikėjimą turiniu. H3: Pasitikėjimas teigiamai veikia turistų ketinimus lankyti. H4: Pasitikėjimas veikia ryšį tarp autentiškumo ir ketinimų. H5: Suderinamumas veikia autentiškumo ir ketinimų ryšį. H6: Turistų įsitraukimas moderoja autentiškumo ir ketinimų ryšį. H7: Žmogaus kurtų ir pažymėtų vizualizacijų autentiškumas didesnis nei DI kurtų ir pažymėtų. H8: Vizualizacijų tipas veikia autentiškumo ir ketinimų ryšį. H9: Nepažymėtų vizualizacijų autentiškumas bus didesnis nei DI pažymėtų.	modeliai, mediacijos ir moderacijos analizės	lankyti atitinkamose vietovėse. Jei pažymėtos, žmogaus kurtos nuotraukos laikomos autentiškesnėmis nei DI, tačiau jei DI autorystė nėra nurodoma, DI sugeneruotos nuotraukos taip pat laikomos gana autentiškomis.	industrijoms, ne tik turizmui ir ilgalaikius DI sprendimų populiarėjimo padarinius.
--	----------------------------------	--	--	--	--	---

Carney et al., 2024	Vartotojų įsitraukimas į DI sukurtą ir pažymėtą turinį socialiniame tinkle „TikTok“	Parasocialinių santykių teorija	Tikslas: Įvertinti, kaip DI sukurtų įrašų žymėjimas veikia vartotojų įsitraukimą socialiniuose tinkluose. Hipotezės: H1: Įrašai, kuriuose yra DI sukurtos medijos žymėjimas, sulauks mažesnio vartotojų įsitraukimo nei tie, kuriuose tokio žymėjimo nėra. H2: Įsitraukimo sumažėjimas nėra susijęs su tikrais ar suvoktais turinio kokybės skirtumais ar su vartotojų susirūpinimu dėl galimo turinio klaidinimo. H3: DI turinio žymėjimas mažina vartotojų įsitraukimą, nes silpnina jų jaučiamą ryšį su turinio kūrėju.	„TikTok“ duomenų analizė (1,1 mln. įrašų, 8 650 naudotojų) bei eksperimentai su dalyviais, kuriuose buvo pateikti DI pažymėti ir nepažymėti įrašai	Tyrimo metu nustatyta, jog vartotojų įsitraukimas į DI pažymėtą turinį mažesnis nei į nepažymėtą turinį. Anot autorių, tai nėra susiję su suvokiamu turinio kokybės ar autentiškumo skirtumu, mat pagrindinė priežastis – sumažėjęs socialinis ryšys su kūrėju. Pozityvus DI žymėjimas („Pažiūrėkite, kokią įspūdingą efektą sukūrė DI!“) taip pat nepadidino įsitraukimo.	Tyrimo apie tolesnes kryptis nėra užsimenama.
Chen et al., 2025	DI atskleidimo poveikis vartotojų įsitraukimo ketinimams trumpųjų vaizdo įrašų atžvilgiu	Heuristinis-sisteminis modelis (HSM)	Tikslas: Ištirti, kaip DI atskleidimas paveikia vartotojų įsitraukimo ketinimus trumpųjų vaizdo įrašų platformose. Hipotezės: H1: DI autorystės atskleidimas turi teigiamą poveikį vartotojų įsitraukimo ketinimams. H2: DI autorystės atskleidimas neigiamai veikia vartotojų suvokiamą turinio kokybę. H3: Suvokiama turinio kokybė teigiamai veikia vartotojų įsitraukimo	Eksperimentinis tyrimas su 479 dalyviais, atliktas internetinėje platformoje „Credamo“. Dalyviai buvo atsitiktinai paskirstyti į grupes atskleidžiant DI autorystę ir jos neatskleidžiant. Duomenys analizuoti naudojant PLS-SEM modelį.	DI autorystės atskleidimas skatina vartotojų įsitraukimą, tačiau tuo pačiu neigiamai veikia jų nuomonę apie turinio kokybę. Kai vartotojai suvokia DI gebėjimus kaip teigiamus (pavyzdžiui, patikimas & lankstus), neigiamas poveikis turinio kokybės vertinimui sumažėja, o jų įsitraukimas didėja.	Autoriai skatina tirti, kaip vartotojai reaguoja į skirtingą DI ir žmogaus bendradarbiavimą. Taip pat siūloma tyrinėti kitus veiksnius, pavyzdžiui, empatiją ar pasitikėjimą DI.

			ketinimus. H4: Vartotojų suvokiami DI gebėjimai moderoja ryšį tarp DI atskleidimo ir suvokiamos turinio kokybės.			
--	--	--	---	--	--	--

3 priedas. Naudotų matavimo skalių ir jų adaptacijų palyginimas tyrime

Matavimo tikslas	Autorius(-iai)	Matavimo skalės turinys originalo kalba	Autorių naudota skalė	Tyrime naudojamas matavimo skalės turinys	Tyrime naudojama skalė
Papildomas klausimas	Chen et al., 2025	Based on the textual description in the scenario, this video was generated by AI.	Septynių balų Likerto skalė (1 = visiškai nesutinku, 7 = visiškai sutinku)	Įvertinkite šį teiginį. Remiantis paveikslėlyje pavaizduota etikete, šį įrašą sukūrė DI.	Septynių balų Likerto skalė (1 = visiškai nesutinku, 7 = visiškai sutinku)
				Argumentuokite savo atsakymą.	Atviras klausimas
Šaltinio ir turinio atitikimas	Lee et al., 2025	I believe there is a contradiction between the advertising message and the source that created the advertisement. I believe the advertising message is consistent with the source that created the advertisement. I believe the source that created the advertisement effectively communicated the intended message.	Penkių balų Likerto skalė (5 = visiškai sutinku, 1 = visiškai nesutinku)	Manau, kad socialinių tinklų įrašas ir jo autorius atitinka vienas <i>kitą</i> . <i>Manau, kad</i> socialinių tinklų įrašo autorius veiksmingai perteikė numatytą žinutę.	Septynių balų Likerto skalė (1 = visiškai nesutinku, 7 = visiškai sutinku)
	Choi, 2024	How relevant is this AI influencer to the product? How congruent is the AI influencer with the product? Overall, how congruent is the AI influencer with the product?	Septynių balų Likerto skalė (1 = visiškai nesutinku, 7 = visiškai sutinku)	Manau, kad šis socialinių tinklų įrašas yra susijęs su jo autoriumi.	
Šaltinio patikimumas	McCroskey & Teven, 1999	Instructions: Please indicate your impression of the person noted below by circling the appropriate number between the pairs of adjectives below. The closer the number is to an adjective, the more certain you are of your evaluation.	Septynių balų semantinio diferencialo skalė	Instrukcijos: Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš	Septynių balų semantinio diferencialo skalė

		Competence Intelligent – Unintelligent Untrained – Trained Inexpert – Expert Informed – Uninformed Incompetent – Competent Bright – Stupid Goodwill Cares about me – Doesn't care about me Has my interests at heart – Doesn't have my interests at heart Self-centered – Not self-centered Concerned with me – Unconcerned with me Insensitive – Sensitive Not understanding – Understanding Trustworthiness Honest – Dishonest Untrustworthy – Trustworthy Honorable – Dishonorable Moral – Immoral Unethical – Ethical Phoney – Genuine		būdvardžių, tuo labiau esate tikras(-a) dėl savo vertinimo. Kompetencija Protingas – Neprotingas Neapmokytas – Apmokytas Ne ekspertas – Ekspertas Informuotas – Neinformuotas Nekompetentingas – Kompetentingas Sumanus – Kvailas Palankumas Rūpinasi manimi – Nesirūpina manimi Rūpinasi mano interesais – Nesirūpina mano interesais Savanaudis – Nesavanaudis Domisi manimi – Nesidomi manimi Nejautrus – Jautrus Nesupratingas – Supratingas Pasitikėjimas Sąžiningas – Nesažiningas Nekeliantis pasitikėjimo – Keliantis pasitikėjimą Garbingas – Negarbingas Moralus – Amoralus Neetiškas – Etiškas Netikras – Tikras	
Turinio patikimumas	Brüns & Meißner, 2024	Posts by [brand] appear trustworthy Posts by [brand] appear honest Posts by [brand] appear convincing	Septynių balų Likerto skalė (1 = visiškai nesutinku, 7 = visiškai sutinku)	Šis socialinių tinklų įrašas kelia pasitikėjimą Šis socialinių tinklų įrašas atrodo įtikinamas Šis socialinių tinklų įrašas atrodo sąžiningas	Septynių balų Likerto skalė (1 = visiškai nesutinku, 7 = visiškai sutinku)
Įsitraukimo ketinimai	Onofrei et al., 2022	After reading the post shared by people in my social media network, I will press the 'like' button	Penkių balų Likerto skalė (5 = visiškai sutinku, 1	Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, paspausčiau „patinka“ po šiuo įrašu	Septynių balų Likerto skalė (1 = visiškai

		After reading the post, I will comment on it After reading the post, I will share it with my friends	= visiškai nesutinku)	Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, palikčiau komentarą po šiuo įrašu4 Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, pasidalinčiau juo su kitais	nesutinku, 7 = visiškai sutinku)
Socialinių tinklų naudojimas	Jia, 2025	Social media usage intensity	‘Mild = About 1 h per day or less’ to ‘Severe = About 5 h per day or more’	Kiek kartų per dieną vidutiniškai naudojate socialiniais tinklais? ○ 1 valandą ir mažiau ○ 2 valandas ○ 3 valandas ○ 4 valandas ○ 5 valandas ir daugiau	
				Kurias socialinių tinklų platformas naudojate dažniausiai? (Galimi keli pasirinkimai) ☐ Facebook ☐ Instagram ☐ TikTok ☐ YouTube ☐ LinkedIn ☐ X (Twitter) ☐ Kita (įrašykite)	
Amžius				Koks yra jūsų amžius?	Atviras klausimas.
Lytis				Jūs esate: ○ Vyras ○ Moteris ○ Kita	
Išsilavinimas				Koks yra aukščiausias jūsų įgytas išsilavinimas? ○ Pradinis išsilavinimas ○ Pagrindinis išsilavinimas	

		☐ Vidurinis išsilavinimas ☐ Aukštesnysis arba specialusis vidurinis išsilavinimas ☐ Aukštasis išsilavinimas	
Pajamos		Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)? ☐ Mažiau nei 500 Eur ☐ 500–999 Eur ☐ 1000–1499 Eur ☐ 1500–1999 Eur ☐ 2000–2999 Eur ☐ 3000 Eur ar daugiau ☐ Nenoriu atsakyti	
Dėmesio patikrinimo klausimas (angl. attention check)		Nurodykite, kokia inovacija buvo pristatoma Jums pateiktame socialinių tinklų įrašė: ☐ Maisto produktas ☐ DI sistema depresijai atpažinti ☐ Biuro kėdė	
Manipuliacijos patikrinimo klausimas (angl. manipulation check)		Įrašas buvo pažymėtas dirbtinio intelekto autorystę nurodančia etikete „AI info“ (liet. dirbtinio intelekto informacija) ☐ Taip ☐ Ne ☐ Nežinau	

4 priedas. Tyrimo anketa

Gerbiamas respondente, Esu Kristina Grybaitė, Kauno technologijos universiteto Marketingo valdymo magistrantūros studijų programos studentė. Atlieku tyrimą, kuriuo siekiu suprasti, kaip įvairūs rinkodaros įrašai socialiniuose tinkluose veikia vartotojų reakcijas.

Apklausos metu Jums bus pateiktas socialinių tinklų įrašo pavyzdys, o po jo – keli trumpi klausimai apie tai, kaip Jūs vertinate įrašo turinį ir jo autorių.

Jūsų dalyvavimas apklausoje yra savanoriškas ir anonimiškas, o visi pateikti duomenys – konfidencialūs. Surinkta informacija bus naudojama tik baigiamojo magistro darbo rengimo tikslais. Anketos pildymas truks apie 6–9 minutes. Jei turėtumėte klausimų arba norėtumėte susipažinti su tyrimo rezultatais, kviečiu susisiekti el. paštu: kristina.grybaite@ktu.lt

1. Ar naudojātės socialiniais tinklais?

- ☐ Taip
- ☐ Ne

Scenarijus 1

KTU Kauno technologijos universitetas/Kaunas University of Technology AI info · 1d. ·

Depresija paliečia milijonus žmonių visame pasaulyje, todėl Kauno technologijos universiteto mokslininkai sukūrė naują dirbtinio intelekto (DI) sprendimą, padedantį ją atpažinti. Ši sistema vertina žmogaus kalbą ir smegenų veiklą – būtent šių dviejų duomenų derinys leidžia tiksliau įvertinti emocinę būklę. 🧠 Tyrimo metu pasiektas net 97 % tikslumas, o balsas pasirinktas kaip itin jautrus emocijų atspindys – jis keičiasi priklausomai nuo nuotaikos, tempo, tono ar energijos. Šis DI sprendimas kuriamas taip, kad būtų užtikrintas žmogaus privatumas, nes nereikia naudoti veido įrašų. 📺 Mokslininkai tikisi, kad ateityje ši technologija prisidės prie greitesnės, tikslesnės ir labiau prieinamos psichikos sveikatos diagnostikos.



👍❤️ 78 žmonės

Patinka Komentuoti Bendrinti

2. Remiantis paveikslėlyje pavaizduota etikete, šį įrašą sukūrė dirbtinis intelektas.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

Argumentuokite savo atsakymą:

Scenarijus 2

KTU Kauno technologijos universitetas/Kaunas University of Technology 1d. · 🌐

🔴 Depresija paliečia milijonus žmonių visame pasaulyje, todėl Kauno technologijos universiteto mokslininkai sukūrė naują dirbtinio intelekto (DI) sprendimą, padedantį ją atpažinti. Ši sistema vertina žmogaus kalbą ir smegenų veiklą – būtent šių dviejų duomenų derinys leidžia tiksliau įvertinti emocinę būklę. 🎧 Tyrimo metu pasiektas net 97 % tikslumas, o balsas pasirinktas kaip itin jautrus emocijų atspindys – jis keičiasi priklausomai nuo nuotaikos, tempo, tono ar energijos. Šis DI sprendimas kuriamas taip, kad būtų užtikrintas žmogaus privatumas, nes nereikia naudoti veido įrašų. 🇸🇮 Mokslininkai tikisi, kad ateityje ši technologija prisidės prie greitesnės, tikslesnės ir labiau prieinamos psichikos sveikatos diagnostikos.



👍❤️ 78 žmonės

👍 Patinka 💬 Komentuoti ➦ Bendrinti

2. Remiantis paveikslėlyje pavaizduota etikete, šį įrašą sukūrė dirbtinis intelektas.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

Argumentuokite savo atsakymą:

Scenarijus 3

KTU Kauno technologijos universitetas/Kaunas University of Technology AI info · 1d. ·

💧 Vandens trūkumas – dažna problema tarp senyvo amžiaus žmonių, ypač turinčių rijimo sunkumų. Kauno technologijos universiteto mokslininkai kuria naują, Lietuvoje analogų neturintį maisto produktą, skirtą sumažinti dehidratacijos ir užspringimo riziką. Tai vieno kąsnio dydžio, ryškių spalvų, gaivaus skonio ir kvapo užkandis, kuris sunkiai suardomas rankomis, bet lengvai tirpsta burnoje. 🍷 Sudarytas net iš 95 % vandens, jis papildytas vitaminais (D, C, B9, B12) ir svarbiais mineralais (geležimi, seleno, cinku), todėl padeda užtikrinti senjorų mitybos poreikius. Šiuo metu vyksta produkto tobulinimas, o netrukus jis bus bandomas ir slaugos namuose.



👍❤️ 78 žmonės

👍 Patinka 💬 Komentuoti ➦ Bendrinti

2. Remiantis paveikslėlyje pavaizduota etikete, šį įrašą sukūrė dirbtinis intelektas.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

Argumentuokite savo atsakymą:

Scenarijus 4



2. Remiantis paveikslėlyje pavaizduota etikete, šį įrašą sukūrė dirbtinis intelektas.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

Argumentuokite savo atsakymą:

3. Įvertinkite kiekvieną iš šių teiginių

Manau, kad socialinių tinklų įrašas ir jo autorius atitinka vienas kitą.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

☐	☐	☐	☐	☐	☐	☐
-----------------------	-----------------------	-----------------------	-----------------------	-----------------------	-----------------------	-----------------------

Manau, kad socialinių tinklų įrašo autorius veiksmingai perteikė numatytą žinutę.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

Manau, kad šis socialinių tinklų įrašas yra susijęs su jo autoriumi.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

4. Įvertinkite kiekvieną iš šių teiginių

Šis socialinių tinklų įrašas kelia pasitikėjimą.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

Šis socialinių tinklų įrašas atrodo sąžiningas.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

Šis socialinių tinklų įrašas atrodo įtikinamas.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

5. Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau esate tikras(-a) dėl savo vertinimo.

Protingas	☐	☐	☐	☐	☐	☐	☐	Neprotingas
Neapmokytas	☐	☐	☐	☐	☐	☐	☐	Apmokytas
Ne ekspertas	☐	☐	☐	☐	☐	☐	☐	Ekspertas
Informuotas	☐	☐	☐	☐	☐	☐	☐	Neinformuotas
Nekompetentingas	☐	☐	☐	☐	☐	☐	☐	Kompetentingas
Sumanus	☐	☐	☐	☐	☐	☐	☐	Kvailas

6. Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau esate tikras(-a) dėl savo vertinimo.

Rūpinasi manimi	☐	☐	☐	☐	☐	☐	☐	Nesirūpina manimi
Rūpinasi mano interesais	☐	☐	☐	☐	☐	☐	☐	Nesirūpina mano interesais

Savanaudis	☐	☐	☐	☐	☐	☐	☐	Nesavanaudis
Domisi manimi	☐	☐	☐	☐	☐	☐	☐	Nesidomi manimi
Nejautrus	☐	☐	☐	☐	☐	☐	☐	Jautrus
Nesupratingas	☐	☐	☐	☐	☐	☐	☐	Supratingas

7. Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau esate tikras(-a) dėl savo vertinimo.

Sąžiningas	☐	☐	☐	☐	☐	☐	☐	Nesąžiningas
Nekeliantis pasitikėjimo	☐	☐	☐	☐	☐	☐	☐	Keliantis pasitikėjimą
Garbingas	☐	☐	☐	☐	☐	☐	☐	Negarbingas
Moralus	☐	☐	☐	☐	☐	☐	☐	Amoralus
Neetiškas	☐	☐	☐	☐	☐	☐	☐	Etiškas
Netikras	☐	☐	☐	☐	☐	☐	☐	Tikras

8. Įvertinkite kiekvieną iš šių teiginių

Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, paspausčiau „patinka“ po šiuo įrašu.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

Palikčiau komentarą po šiuo įrašu.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

Pasidalinčiau šiuo įrašu su kitais.

Visiškai nesutinku	Nesutinku	Labiau nesutinku	Nei sutinku, nei nesutinku	Labiau sutinku	Sutinku	Visiškai sutinku
☐	☐	☐	☐	☐	☐	☐

9. Kiek laiko vidutiniškai per dieną naudojate socialiniais tinklais?

- ☐ 1 valandą ar mažiau
- ☐ 2 valandas
- ☐ 3 valandas
- ☐ 4 valandas
- ☐ 5 valandas ar daugiau

10. Kurias socialinių tinklų platformas naudojate dažniausiai? (Galite pasirinkti kelias)

- ☐ Facebook
 - ☐ Instagram
 - ☐ TikTok
 - ☐ YouTube
 - ☐ LinkedIn
 - ☐ X (Twitter)
 - ☐ Kita (įrašykite): _____
-

11. Jūsų amžius (įrašykite)

12. Jūs esate

- ☐ Vyras
- ☐ Moteris
- ☐ Kita

13. Koks yra aukščiausias jūsų įgytas išsilavinimas?

- ☐ Pradinis išsilavinimas
- ☐ Pagrindinis išsilavinimas
- ☐ Vidurinis išsilavinimas
- ☐ Aukštesnysis arba specialusis vidurinis išsilavinimas
- ☐ Aukštasis išsilavinimas

14. Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?

- ☐ Mažiau nei 500 Eur
 - ☐ 500–999 Eur
 - ☐ 1000–1499 Eur
 - ☐ 1500–1999 Eur
 - ☐ 2000–2499 Eur
 - ☐ 2500 Eur ar daugiau
 - ☐ Nenoriu atsakyti
-

15. Nurodykite, kokia inovacija buvo pristatoma Jums pateiktame socialinių tinklų įrašė:

- ☐ Maisto produktas
- ☐ DI sistema depresijai atpažinti
- ☐ Biuro kėdė

16. Įrašas buvo pažymėtas dirbtinio intelekto autorystę nurodančia etikete „AI info“ (liet. dirbtinio intelekto informacija)

- ☐ Taip
- ☐ Ne
- ☐ Nežinau

Dėkoju, kad dalyvavote tyrime!

Šio tyrimo tikslas buvo ištirti, kaip socialiniuose tinkluose publikuojamas dirbtinio intelekto (DI) sukurtas reklaminis turinys veikia vartotojų suvokiamą šaltinio ir turinio patikimumą bei įsitraukimo ketinimus.

Apklausos metu Jums buvo pateiktas vienas iš keturių įrašų, kurie visi buvo sukurti dirbtinio intelekto pagalba, tačiau kai kuriuose jų buvo nurodyta, kad įrašas yra sukurtas DI, o kituose – ne. Be to, dalis įrašų buvo apie su DI susijusius produktus, o kita dalis – apie nesusijusias inovacijas. Tokiu būdu buvo siekiama įvertinti, kaip skirtingos žymos ir turinio temos veikia Jūsų požiūrį į įrašą ir jo autorių.

Kad tyrimas būtų kuo objektyvesnis, kai kurie elementai buvo pateikti neatskleidžiant visos informacijos iš karto, tačiau viskas buvo suplanuota laikantis mokslinių tyrimų etikos principų.

Jūsų atsakymai yra labai vertingi ir bus naudojami tik akademiniams tikslais, rengiant magistro baigiamąjį darbą Kauno technologijos universitete.

Jeigu turite klausimų arba norėtumėte sužinoti tyrimo rezultatus, galite susisiekti el. paštu: kristina.grybaite@ktu.lt

5 priedas. Užklausos, pateiktos scenarijų turiniui sukurti

Pateikus 1 užklausą „Sukurk socialinio tinklo įrašą „Facebook“ platformai, kuris būtų informacinio tipo. Įrašo tikslas – aiškiai, suprantamai ir patraukliai pristatyti žemiau pateiktą produktą ar inovaciją. Įrašas turi būti skirtas plačiai auditorijai, todėl rašyk paprastai, bet informatyviai, taisyklinga lietuvių kalba, vengdamas sudėtingų mokslo terminų ar angliškų santrumpų. Jei reikia, apibūdink tokius terminus suprantamai, o trumpinį rašyk skliausteliuose.

Įrašas turi būti parašytas viena pastraipa, iki 150 žodžių, trumpais, aiškiais sakiniais. Tonas – neutralus, informacinis, bet draugiškas, be reklaminių šūkių ar perteklinio emociingumo. Gali naudoti iki 3 jaustukų, paskirstytų skirtingose teksto vietose (ne šalia vienas kito).

Naudok šį turinį teksto kūrimui: Depresija – vienas iš dažniausiai pasitaikančių psichikos sveikatos sutrikimų. Pasaulyje net 280 milijonų žmonių yra paveikti šios diagnozės, todėl Kauno technologijos universiteto (KTU) tyrėjai sukūrė dirbtinio intelekto (DI) modelį, padedantį atpažinti depresiją remiantis tiek kalbos, tiek smegenų veiklos duomenimis.

Šis multimodalinis metodas, jungiantis du skirtingus duomenų šaltinius, leidžia tiksliau ir objektyviau analizuoti emocinę būklę, atverdamas galimybes naujam depresijos diagnostikos etapui.

„Depresija yra vienas labiausiai paplitusių psichikos sutrikimų, sukeliančių skaudžias pasekmes tiek individui, tiek visuomenei, todėl kuriame naują, objektyvesnę diagnostikos būdą, kuris ateityje galėtų tapti prieinamas kiekvienam“, – sako KTU profesorius ir vienas iš išradimo autorių Rytis Maskeliūnas.

Nors daugelis depresijos diagnostikos tyrimų tradiciškai remiasi vienu duomenų tipu, mokslininkų teigimu, naujasis metodas, atsižvelgiantis į kelis aspektus, suteikia daugiau informacijos apie žmogaus emocinę būklę.

Įspūdingas tikslumas, naudojant balso ir smegenų veiklos duomenis

Toks kalbos ir smegenų veiklos duomenų derinys depresijos diagnostikoje pasiekė įspūdingą 97,53 proc. tikslumą, ženkliai viršydamas alternatyvių metodų rezultatus. „Taip yra todėl, kad balsas tyrimą papildė tokiais duomenimis, kurių negalime ir dar nemokame išgauti iš smegenų“, – aiškina R. Maskeliūnas.

Musyyab Yousufi, KTU doktorantas, prisidėjęs prie šio išradimo kūrimo, teigia, jog duomenų pasirinkimas buvo detalai apmąstytas: „Nors galima manyti, jog, pavyzdžiui, veido išraiškos atskleidžia daugiau apie žmogų, balsą pasirinkome, nes jo pokyčiai gali subtiliai atskleisti emocinę būklę, diagnozei paveikiant kalbos tempą, intonaciją ir bendrą energiją, o minėtosios veido išraiškos gali būti labiau kontroliuojamos“.

Be to, tyrėjai suprato, jog retas depresija sergantis pacientas norėtų savo veido įrašų dalytis viešai. Priešingai nei smegenų elektrinės veiklos (EEG) ar balso duomenys, veidas – tiesiogiai asmenį identifikuojanti informacija.

„Negalime pažeisti pacientų privatumo, o ir tolimesniam naudojimui mūsų pasirinktų duomenų derinimas yra perspektyvesnis“, – tvirtina KTU informatikos fakulteto (IF) profesorius.

Mokslininkas pabrėžia, jog jie nėra medicinos ekspertai, todėl tyrimo atlinkti tiesiogiai su pacientais negali. Dėl šios priežasties visi duomenys buvo gauti iš MODMA (angl. Multimodal Open Dataset for Mental Disorder Analysis) duomenų rinkinio. EEG duomenys buvo renkami ir įrašinėjami penkias minutes, dalyviams būnant budriems, ramybės būsenoje, užmerkus akis ir nejudant.

Atliekant garso tyrimo eksperimentines užduotis pacientai dalyvavo klausimų ir atsakymų sesijoje bei kelse veiklose, orientuotose į skaitymą ir paveikslėlių apibūdinimą, siekiant užfiksuoti jų natūralią kalbą ir kognityvinę būklę.

DI turės išmokti diagnozę pagrįsti

Surinkti EEG ir garso signalai buvo transformuoti į spektrogramas, leidžiančias duomenis pavaizduoti vizualiai. Kad šiuos duomenis būtų galima palyginti, buvo pritaikyti specialūs filtrai triukšmui pašalinti ir panaudotas modifikuotas „DenseNet-121“ giluminio mokymosi modelis, gebantis atpažinti depresijos požymius vaizduose. Kiekvienas vaizdas atspindėjo signalų pokyčius laikui bėgant, EEG parodė smegenų veiklos bangų formas, o garsas – dažnių ir intensyvumo pasiskirstymą.

Į šį modelį buvo įtrauktas pasirinktinis klasifikavimo sluoksnis, apmokytas skirstyti duomenis į sveikų arba depresija sergančių žmonių klases. Vertinant sėkmingą klasifikaciją ir buvo įvertintas programos tikslumas.

Ateityje šis DI modelis galėtų pagreitinti depresijos diagnostiką ar ją net paversti nuotoline bei sumažinti subjektyvių vertinimų riziką. Tam reikalingi tolimesni klinikiniai tyrimai bei programos tobulinimai, tačiau R. Maskeliūnas teigia, jog susiduriama ir su iššūkiais.

„Pagrindinė tokių tyrimų problema – duomenų trūkumas, nes žmonės nenori atsiverti, dalytis“, – teigia jis.

Kitas svarbus Programų inžinerijos katedros profesoriaus paminėtas aspektas yra tai, jog nepakanka vien gero tikslumo sukurtam DI algoritmui, nes gydytojas turi suprasti, kas lėmė tokį diagnostinį rezultatą. „Algoritmas dar turi išmokyti suprantamai paaiškinti diagnozę“, – šypsosi mokslininkas.

R. Maskeliūno teigimu, su šia problema susiduriama ne vienoje srityje dėl didėjančio poreikio taikyti DI ten, kur priimami sprendimai tiesiogiai veikia žmones, pavyzdžiui, sveikatos priežiūroje, finansuose ir teisinėje sistemoje.

Todėl dabar populiarėja XAI (angl. explainable artificial intelligence) paaiškinamoji dirbtinio intelekto technologija, siekianti naudotojui atskleisti, kodėl modelis priima konkrečius sprendimus, taip padidinti pasitikėjimą DI, ChatGPT sugeneravo atsakymą:  *Depresija paliečia milijonus žmonių visame pasaulyje, todėl Kauno technologijos universiteto mokslininkai sukūrė naują dirbtinio intelekto (DI) sprendimą, padedantį ją atpažinti. Ši sistema vertina žmogaus kalbą ir smegenų veiklą – būtent šių dviejų duomenų derinys leidžia tiksliau įvertinti emocinę būklę.*  *Tyrimo metu pasiektas net 97 proc. tikslumas, o balsas pasirinktas kaip itin jautrus emocijų atspindys – jis keičiasi priklausomai nuo nuotaikos, tempo, tono ar energijos. Šis DI sprendimas kuriamas taip, kad būtų užtikrintas žmogaus privatumas, nes nereikia naudoti veido įrašų.*  *Mokslininkai tikisi, kad ateityje ši technologija prisidės prie greitesnės, tikslesnės ir labiau prieinamos psichikos sveikatos diagnostikos.*

Pateikus 2 užklausą „Sukurk socialinio tinklo įrašą „Facebook“ platformai, kuris būtų informacinio tipo. Įrašo tikslas – aiškiai, suprantamai ir patraukliai pristatyti žemiau pateiktą produktą ar inovaciją. Įrašas turi būti skirtas plačiai auditorijai, todėl rašyk paprastai, bet informatyviai, taisyklinga lietuvių kalba, vengdamas sudėtingų mokslo terminų ar anglišų santrumpų. Jei reikia, apibūdink tokius terminus suprantamai, o trumpinį rašyk skliausteliuose.

Įrašas turi būti parašytas viena pastraipa, iki 150 žodžių, trumpais, aiškiais sakiniais. Tonas – neutralus, informacinis, bet draugiškas, be reklaminių šūkių ar perteklinio emocingumo. Gali naudoti iki 3 jaustukų, paskirstytų skirtingose teksto vietose (ne šalia vienas kito).

Naudok šį turinį teksto kūrimui: Viena iš pagrindinių problemų, su kuria susiduria daugelis senyvo amžiaus žmonių – skysčių trūkumas arba dehidratacija. Ji gali atsirasti dėl psichologinių ar fizinių sutrikimų, lemiančių pakitusius senjorų gebėjimus savarankiškai pavalgyti ar atsigerti. Kauno technologijos universiteto (KTU) mokslininkai kuria specialų maisto produktą dehidratacijos, o taip pat ir mirtino užspringimo rizikos mažinimui. Netrukus jį išbandyti galės ir patys senjorai.

„Ši inovacija – tai sunkiai pirštais suardoma, bet lengvai burnoje ištirpstančia membrana apdengtas, vieno kąsnio dydžio maisto produktas, išsiskiriantis ryškiomis, pastebimomis spalvomis bei gaišiu skoniu ir aromatu“, – naują apibūdina KTU Cheminės technologijos fakulteto (CTF) profesorė Daiva Leskauskaitė.

Didžiausia tokio naujojo produkto nauda – jo sudėtis. Sudarytas iš 95 proc. vandens, praturtintas D, C, B9 ir B12 vitaminais bei geležies, seleno ir cinko mineralais, šis produktas atitiks rekomenduojamus suvartoti mineralinių medžiagų bei vitaminų kiekius būtent senyvo amžiaus žmonėms.

Šiuo metu maisto produktų kūrimo ir tobulinimo procesas, apie kurį magistro baigiamąjį darbą ruošia KTU ir LSMU bendros magistrantūros studijų programos Medicininė chemija studentė Enrika Lazickaitė, dar tebevyksta, tačiau KTU Maisto mokslo ir technologijos katedros profesorė teigia, jog nemažai jau yra padaryta.

„Rasta daug vandens turinčių maisto produktų struktūrizavimui naudojamų biopolimerų kompozicija, nustatytos galimybės į produktą įdėti vitaminus ir mineralines medžiagas, įvertintas šių komponentų stabilumas produkto technologinio proceso ir laikymo metu bei sukurti rutulio formos vieno kąsnio produktai“, – pasakoja Prof. D. Leskauskaitė.

Laukia labai svarbus tyrimo etapas – nustatyti naujų produktų priimtinumą slaugos namų senyvo amžiaus gyventojų grupėje. Tyrimo rezultatai bus naudojami atrenkant produktų skonį, aromatą bei spalvą.

„Tuo pat metu atliksime į produktus įdėtų vitaminų ir mineralų atpalaidavimo plonajame žarnyne in vitro tyrimus. Gauti rezultatai leis spręsti apie įdėtų mikrokomponentų biopasisavinimo efektyvumą senjorų organizmuose“, – papildo KTU mokslininkė.

Paskutinis etapas – bandomosios partijos gamyba, kurios metu bus įvertinta produktų pakuotė bei atlikti paskutiniai gamybos technologinių parametrų derinimai.

Analogų Lietuvoje neturintis produktas

Kalbėdama apie priežastis, nulemiančias nepakankamą senjorų skysčių vartojimą, D. Leskauskaitė teigia, kad jos įvairios, tačiau viena iš pagrindinių – rijimo sutrikimai (disfagija).

„Hospitalizuotiems pacientams, turintiems rijimo sutrikimą, yra didesnė dehidratacijos rizika nei tiems, kuriems ši liga nediagnozuota. Tokiems žmonėms reikalinga vartoti modifikuotos tekstūros skysčius, tačiau vyresnio amžiaus žmonės dažnai jų nemėgsta, todėl vartoja nepakankamą kiekį“, – teigia KTU Maisto mokslo ir technologijos katedros profesorė.

Tačiau D. Leskauskaitė papildo, kad dehidrataciją senyvame amžiuje gali sukelti ir kitos problemos, tarp kurių – depresija, demencija ar net Alzheimerio liga.

„Pacientai pamiršta atsigerti, pavalgyti, o negalėjimas savarankiškai naudotis stalo įrankiais, atbaido sergančiuosius nuo be kokių pastangų maisto atžvilgiu. Taip atsiranda papildomos problemos“, – teigia KTU mokslininkė.

Pasak prof. D. Leskauskaitės, tiksliniai produkto vartotojai ir nustatytos pagrindinės jų dehidratacijos priežastys privertė daug dėmesio skirti ne tik produkto sudėčiai, bet ir išorinėms savybėms – dydžiui, tekstūrai bei patogiam jo paėmimui.

Mitybos bei geriatrijos specialistai teigia, kad maistas, specialiai paruoštas valgyti pirštais, lengvai suimamas bei transportuojamas iš lėkštės į burną, gali būti naudingas senyvo amžiaus žmonėms su kognityviniais sutrikimais bei tapti dehidrataciją mažinančiu aspektu.

Rijimo sutrikimai pavojingi ne tik tuo, kad senyvo žmogaus organizmas nepakankamai aprūpinamas vandeniu ir maistinėmis medžiagomis, bet ir dėl padidėjusios aspiracinės pneumonijos rizikos. Šiuo metu rinkoje esantys šioms rizikoms mažinti skirti produktai yra įvairūs sutirštinti, modifikuotos tekstūros gėrimai.

Ir nors tai tinkama terapinė strategija, siekianti sumažinti užspringimo riziką disfagija sergantiems pacientams, nemažai tyrimų demonstruoja, kad tik sutirštintų skysčių vartojimas neturi ženklios įtakos vandens kiekio padidėjimui šių pacientų organizme, o ir tokie gėrimai nėra priimtini visiems.

„Jokie specializuoti mitybos papildai, skirti rijimo problemų turintiems asmenims, Lietuvoje neegzistuoja, todėl naujas, tiek naudingumo, tiek patrauklumo atžvilgiu pažangesnis produktas – ypač reikalingas“, – argumentuoja prof. D. Leskauskaitė“, ChatGPT sugeneravo atsakymą:  Vandens trūkumas – dažna problema tarp senyvo amžiaus žmonių, ypač turinčių rijimo sunkumų. Kauno technologijos universiteto mokslininkai kuria naują, Lietuvoje analogų neturintį maisto produktą, skirtą sumažinti dehidratacijos ir užspringimo riziką. Tai vieno kėsnio dydžio, ryškių spalvų, gaivaus skonio ir kvapo užkandis, kuris sunkiai suardomas rankomis, bet lengvai tirpsta burnoje.  Sudarytas net iš 95 proc. vandens, jis papildytas vitaminais (D, C, B9, B12) ir svarbiais mineralais (geležimi, seleno, cinku), todėl padeda užtikrinti senjorų mitybos poreikius. Šiuo metu vyksta produkto tobulinimas, o netrukus jis bus bandomas ir slaugos namuose.

6 priedas. Tyrimo respondentų charakteristikos

Statistics					
		Lytis	Amžius	Išsilavinimas	Pajamos
N	Valid	372	372	372	372
	Missing	0	0	0	0
Mean		1,60	29,44	2,98	3,91
Median		2,00	25,50	3,00	3,00
Mode		2	25	3	3
Std. Deviation		,501	10,076	1,217	1,808

Lytis					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Vyras	150	40,3	40,3	40,3
	Moteris	220	59,1	59,1	99,5
	Kita	2	,5	,5	100,0
	Total	372	100,0	100,0	

Amžius					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	15	2	,5	,5	,5
	18	2	,5	,5	1,1
	20	16	4,3	4,3	5,4
	21	22	5,9	5,9	11,3
	22	27	7,3	7,3	18,5
	23	36	9,7	9,7	28,2
	24	34	9,1	9,1	37,4
	25	47	12,6	12,6	50,0
	26	38	10,2	10,2	60,2
	27	10	2,7	2,7	62,9
	28	12	3,2	3,2	66,1
	29	18	4,8	4,8	71,0
	30	14	3,8	3,8	74,7
	31	4	1,1	1,1	75,8
	32	4	1,1	1,1	76,9
	34	6	1,6	1,6	78,5
	35	6	1,6	1,6	80,1
	36	6	1,6	1,6	81,7
	37	2	,5	,5	82,3
	38	2	,5	,5	82,8
	39	10	2,7	2,7	85,5
	40	4	1,1	1,1	86,6
	41	2	,5	,5	87,1
	42	2	,5	,5	87,6
	43	2	,5	,5	88,2
	45	2	,5	,5	88,7
	46	4	1,1	1,1	89,8
	47	4	1,1	1,1	90,9
	49	6	1,6	1,6	92,5
	50	2	,5	,5	93,0
	51	4	1,1	1,1	94,1
	52	6	1,6	1,6	95,7
	53	2	,5	,5	96,2
	54	2	,5	,5	96,8
	59	4	1,1	1,1	97,8
	60	2	,5	,5	98,4
	61	4	1,1	1,1	99,5
	67	2	,5	,5	100,0
	Total	372	100,0	100,0	

Statistics

Koks yra jūsų amžius?

N	Valid	372
	Missing	0
Mean		29,44
Median		25,50
Mode		25
Std. Deviation		10,076
Percentiles	25	23,00
	50	25,50
	75	31,00

Amžius_Kvart

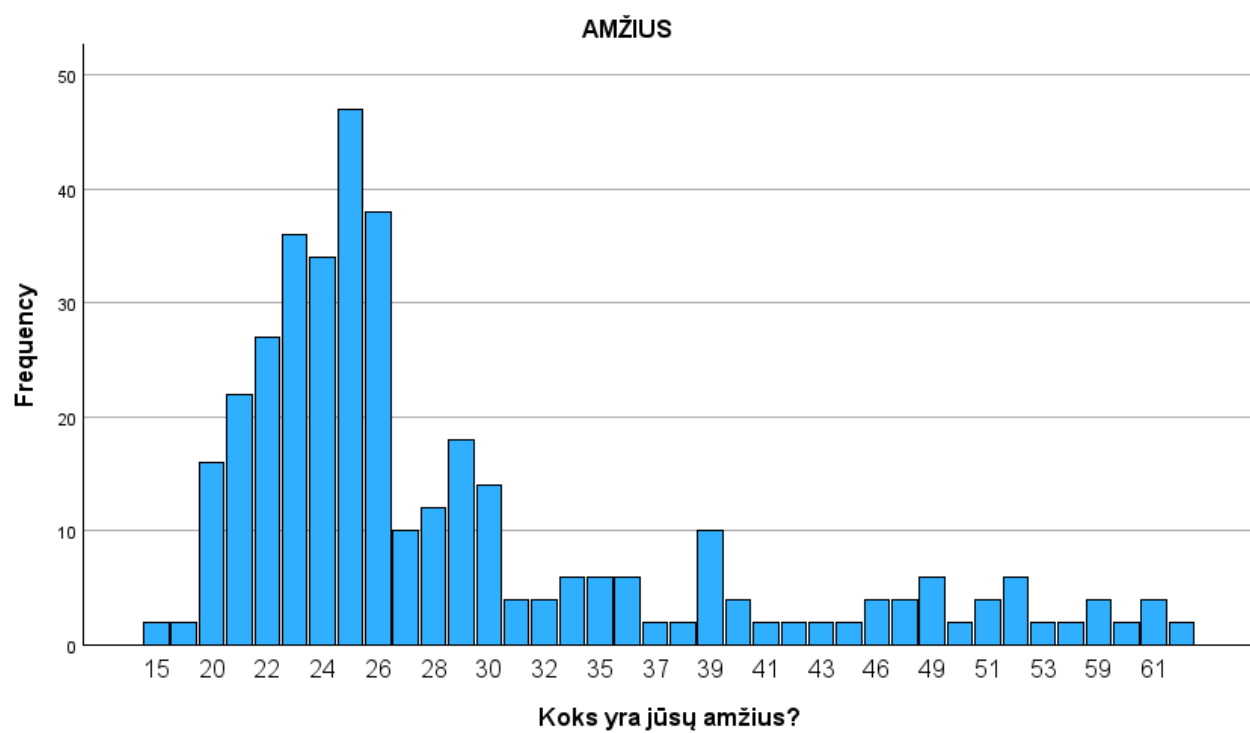
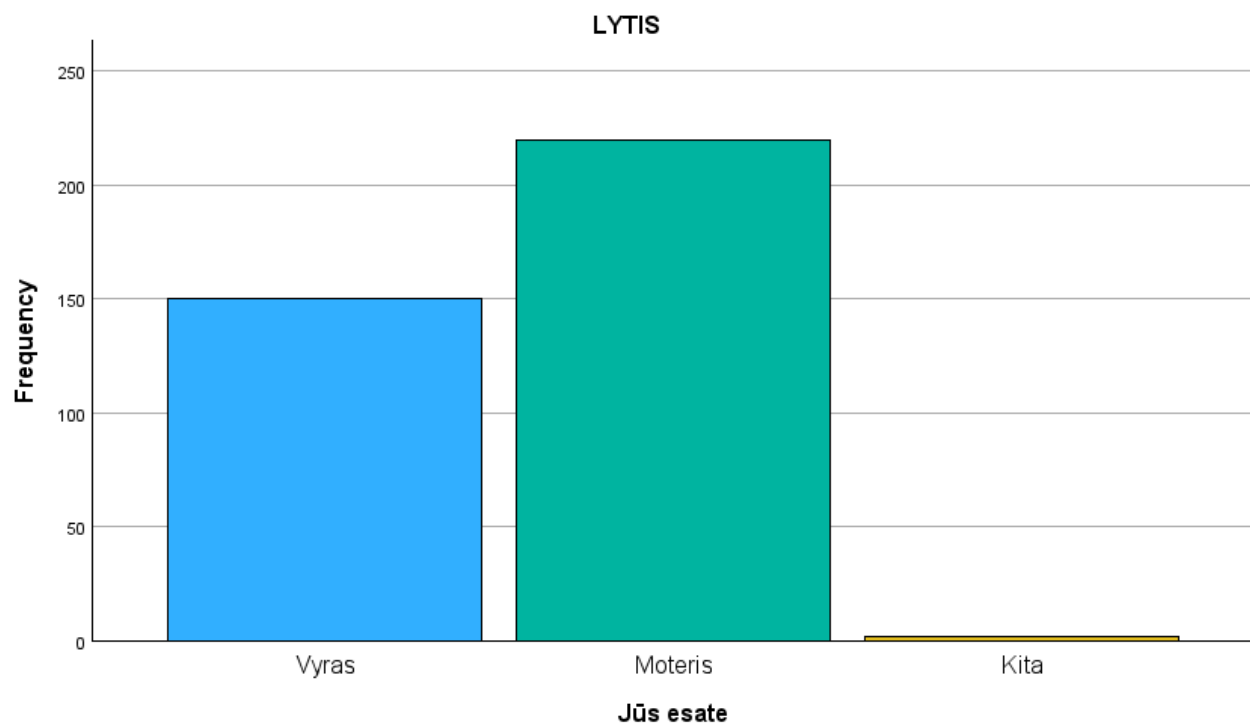
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	15–23 m.	105	28,2	28,2	28,2
	24–25 m.	81	21,8	21,8	50,0
	26–31 m.	96	25,8	25,8	75,8
	32–67 m.	90	24,2	24,2	100,0
	Total	372	100,0	100,0	

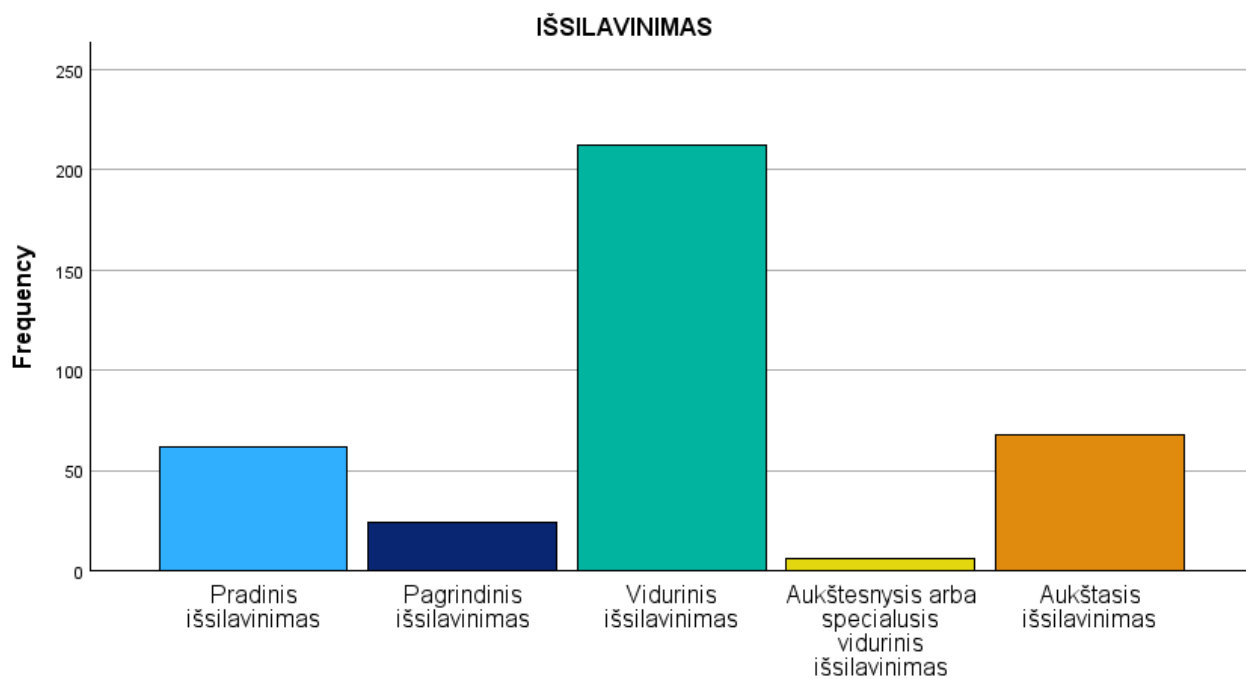
Išsilavinimas

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Pradinis išsilavinimas	62	16,7	16,7	16,7
	Pagrindinis išsilavinimas	24	6,5	6,5	23,1
	Vidurinis išsilavinimas	212	57,0	57,0	80,1
	Aukštesnysis arba specialusis vidurinis išsilavinimas	6	1,6	1,6	81,7
	Aukštasis išsilavinimas	68	18,3	18,3	100,0
	Total	372	100,0	100,0	

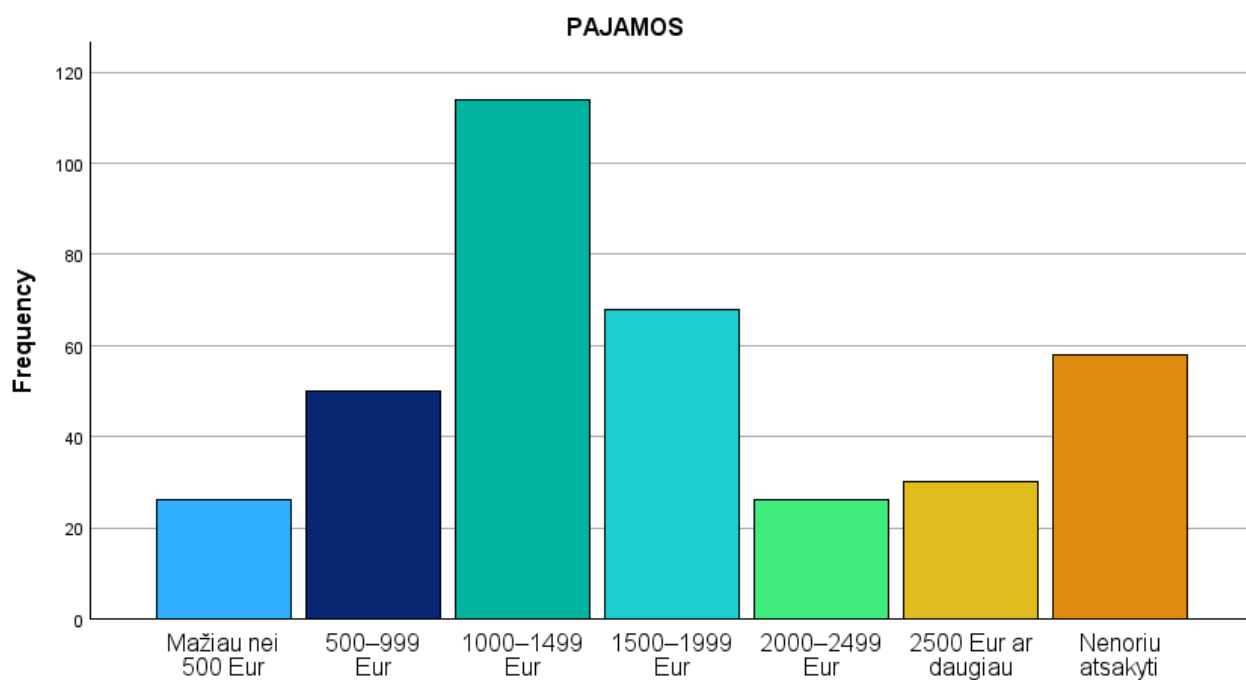
Pajamos

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Mažiau nei 500 Eur	26	7,0	7,0	7,0
	500–999 Eur	50	13,4	13,4	20,4
	1000–1499 Eur	114	30,6	30,6	51,1
	1500–1999 Eur	68	18,3	18,3	69,4
	2000–2499 Eur	26	7,0	7,0	76,3
	2500 Eur ar daugiau	30	8,1	8,1	84,4
	Nenorių atsakyti	58	15,6	15,6	100,0
	Total	372	100,0	100,0	





Koks yra aukščiausias jūsų įgytas išsilavinimas?

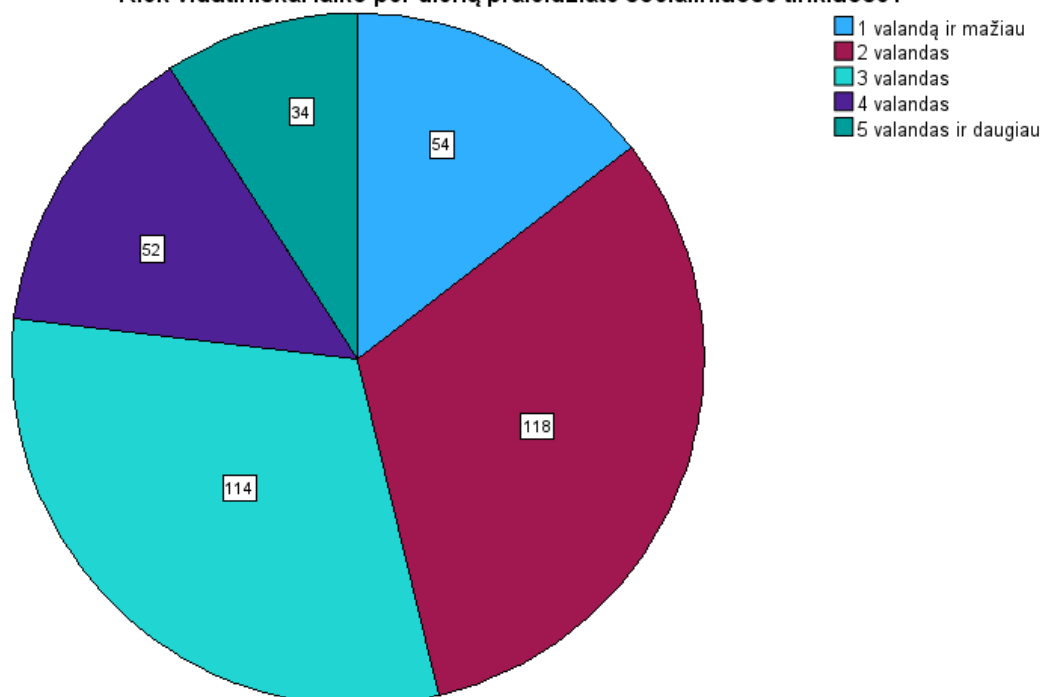


Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?

Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose?

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	1 valandą ir mažiau	54	14,5	14,5	14,5
	2 valandas	118	31,7	31,7	46,2
	3 valandas	114	30,6	30,6	76,9
	4 valandas	52	14,0	14,0	90,9
	5 valandas ir daugiau	34	9,1	9,1	100,0
	Total	372	100,0	100,0	

Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose?



Facebook

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ne	84	22,6	22,6	22,6
	Taip	288	77,4	77,4	100,0
	Total	372	100,0	100,0	

Instagram

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ne	76	20,4	20,4	20,4
	Taip	296	79,6	79,6	100,0
	Total	372	100,0	100,0	

TikTok

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ne	244	65,6	65,6	65,6
	Taip	128	34,4	34,4	100,0
	Total	372	100,0	100,0	

YouTube

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ne	160	43,0	43,0	43,0
	Taip	212	57,0	57,0	100,0
	Total	372	100,0	100,0	

LinkedIn

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ne	290	78,0	78,0	78,0
	Taip	82	22,0	22,0	100,0
	Total	372	100,0	100,0	

X (Twitter)

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ne	350	94,1	94,1	94,1
	Taip	22	5,9	5,9	100,0
	Total	372	100,0	100,0	

SocTinklai_Spotify

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ne	368	98,9	98,9	98,9
	Taip	4	1,1	1,1	100,0
	Total	372	100,0	100,0	

SocTinklai_Reddit

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ne	358	96,2	96,2	96,2
	Taip	14	3,8	3,8	100,0
	Total	372	100,0	100,0	

SocTinklai_Pinterest

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ne	368	98,9	98,9	98,9
	Taip	4	1,1	1,1	100,0
	Total	372	100,0	100,0	

SocTinklai_Mastodon

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ne	370	99,5	99,5	99,5
	Taip	2	,5	,5	100,0
	Total	372	100,0	100,0	

SocTinklai_Discord

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Ne	370	99,5	99,5	99,5
	Taip	2	,5	,5	100,0
	Total	372	100,0	100,0	

7 priedas. Faktorinės analizės rezultatai. Suvokiamas šaltinio patikimumas

KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		,909
Bartlett's Test of Sphericity	Approx. Chi-Square	3667,631
	df	105
	Sig.	<,001

Communalities

	Initial	Extraction
Špasitik_1_rev	1,000	,801
Špasitik_3_rev	1,000	,826
Špasitik_4_rev	1,000	,810
Škompet_1_rev	1,000	,706
Škompet_2	1,000	,765
Škompet_3	1,000	,704
Škompet_4_rev	1,000	,597
Škompet_5	1,000	,796
Škompet_6_rev	1,000	,704
Špalank_1_rev	1,000	,792
Špalank_2_rev	1,000	,762
Špalank_3	1,000	,753
Špalank_4_rev	1,000	,771
Špalank_5	1,000	,754
Špalank_6	1,000	,730

Extraction Method: Principal Component Analysis.

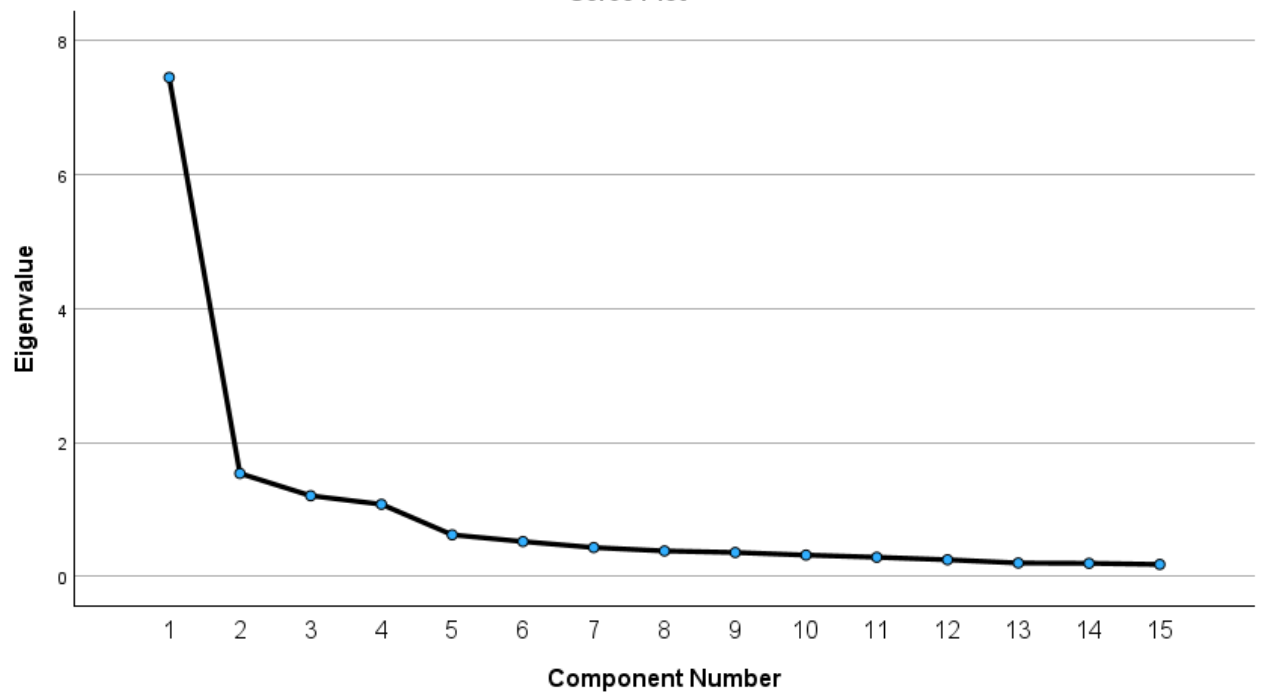
Total Variance Explained

Component	Total	Initial Eigenvalues		Extraction Sums of Squared Loadings			Rotation Sums of Squared Loadings ^a Total
		% of Variance	Cumulative %	Total	% of Variance	Cumulative %	
1	7,452	49,679	49,679	7,452	49,679	49,679	5,420
2	1,538	10,250	59,929	1,538	10,250	59,929	4,229
3	1,204	8,027	67,956	1,204	8,027	67,956	3,510
4	1,076	7,176	75,132	1,076	7,176	75,132	5,076
5	,620	4,135	79,267				
6	,520	3,465	82,732				
7	,430	2,869	85,600				
8	,380	2,530	88,131				
9	,356	2,376	90,507				
10	,318	2,119	92,626				
11	,285	1,903	94,529				
12	,248	1,653	96,182				
13	,200	1,331	97,513				
14	,195	1,299	98,812				
15	,178	1,188	100,000				

Extraction Method: Principal Component Analysis.

a. When components are correlated, sums of squared loadings cannot be added to obtain a total variance.

Scree Plot



Pattern Matrix^a

	Component			
	1	2	3	4
Škompet_3	,854			
Škompet_2	,791			
Škompet_5	,765			
Škompet_1_rev	,624			
Škompet_6_rev	,576			
Škompet_4_rev	,474			
Špalank_4_rev		,908		
Špalank_1_rev		,859		
Špalank_2_rev		,824		
Špalank_3			,775	
Špalank_5			,730	
Špalank_6			,676	
Špasitik_3_rev				-,851
Špasitik_1_rev				-,840
Špasitik_4_rev				-,797

Extraction Method: Principal Component Analysis.

Rotation Method: Oblimin with Kaiser Normalization.

a. Rotation converged in 10 iterations.

8 priedas. Faktorinės analizės rezultatai. Vartotojų įsitraukimo ketinimai

Correlation Matrix

		[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, paspausčiau „patinka“ po šiuo įrašu] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, palikčiau komentarą po šiuo įrašu] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, pasidalinčiau juo su kitais] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.
Correlation	[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, paspausčiau „patinka“ po šiuo įrašu] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,000	,508	,659
	[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, palikčiau komentarą po šiuo įrašu] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,508	1,000	,677
	[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, pasidalinčiau juo su kitais] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,659	,677	1,000

KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		,681
Bartlett's Test of Sphericity	Approx. Chi-Square	441,691
	df	3
	Sig.	<,001

Communalities

	Initial	Extraction
[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, paspausčiau „patinka“ po šiuo įrašu] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,000	,695
[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, palikčiau komentarą po šiuo įrašu] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,000	,711
[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, pasidalinčiau juo su kitais] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,000	,827

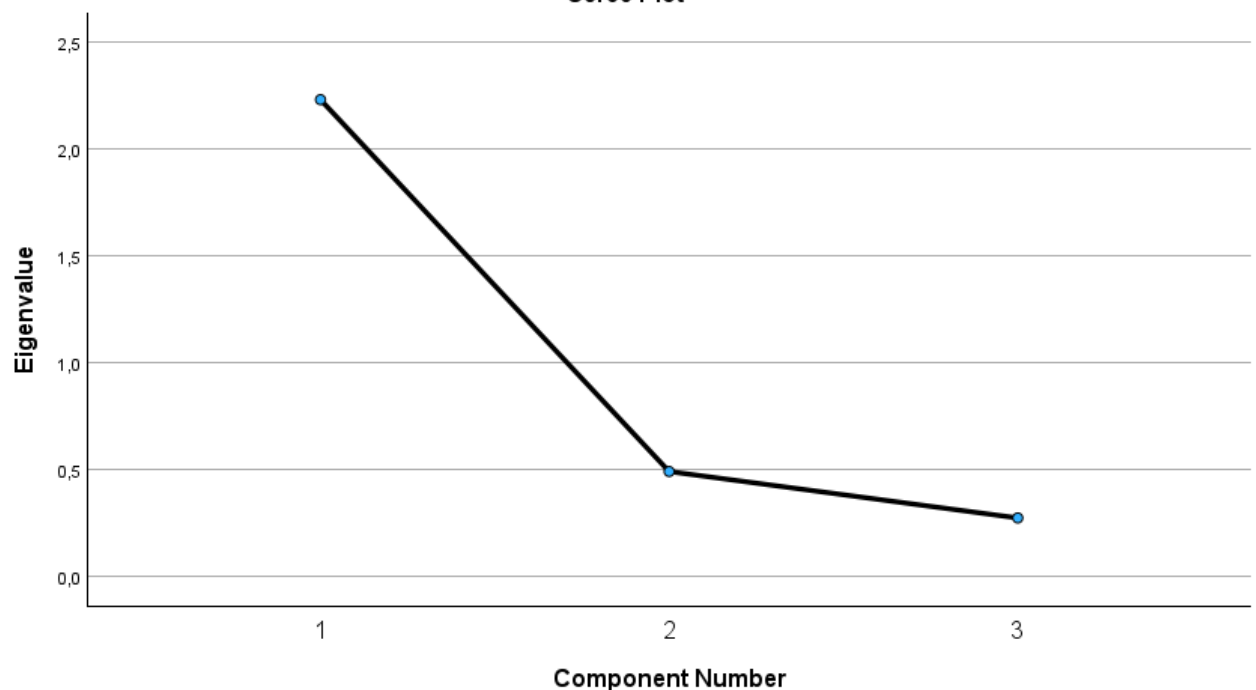
Extraction Method: Principal Component Analysis.

Total Variance Explained

Component	Total	Initial Eigenvalues		Extraction Sums of Squared Loadings		
		% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	2,233	74,422	74,422	2,233	74,422	74,422
2	,492	16,412	90,835			
3	,275	9,165	100,000			

Extraction Method: Principal Component Analysis.

Scree Plot



Component Matrix^a

	Component 1
[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, pasidalinčiau juo su kitais] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,909
[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, palikčiau komentarą po šiuo įrašu] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,843
[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, paspausčiau „patinka“ po šiuo įrašu] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,834

Extraction Method: Principal Component Analysis.

a. 1 components extracted.

9 priedas. Faktorinės analizės rezultatai. Suvokiamas šaltinio ir turinio atitikimas

Correlation Matrix

		[Manau, kad šis socialinių tinklų įrašas ir jo autorius atitinka vienas kitą.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	[Manau, kad socialinių tinklų įrašo autorius veiksmingai perteikė numatytą žinutę.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	[Manau, kad šis socialinių tinklų įrašas yra susijęs su jo autoriumi.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.
Correlation	[Manau, kad šis socialinių tinklų įrašas ir jo autorius atitinka vienas kitą.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,000	,475	,607
	[Manau, kad socialinių tinklų įrašo autorius veiksmingai perteikė numatytą žinutę.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,475	1,000	,316
	[Manau, kad šis socialinių tinklų įrašas yra susijęs su jo autoriumi.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,607	,316	1,000

KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		,611
Bartlett's Test of Sphericity	Approx. Chi-Square	265,140
	df	3
	Sig.	<,001

Communalities

	Initial	Extraction
[Manau, kad šis socialinių tinklų įrašas ir jo autorius atitinka vienas kitą.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,000	,778
[Manau, kad socialinių tinklų įrašo autorius veiksmingai perteikė numatytą žinutę.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,000	,512
[Manau, kad šis socialinių tinklų įrašas yra susijęs su jo autoriumi.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,000	,653

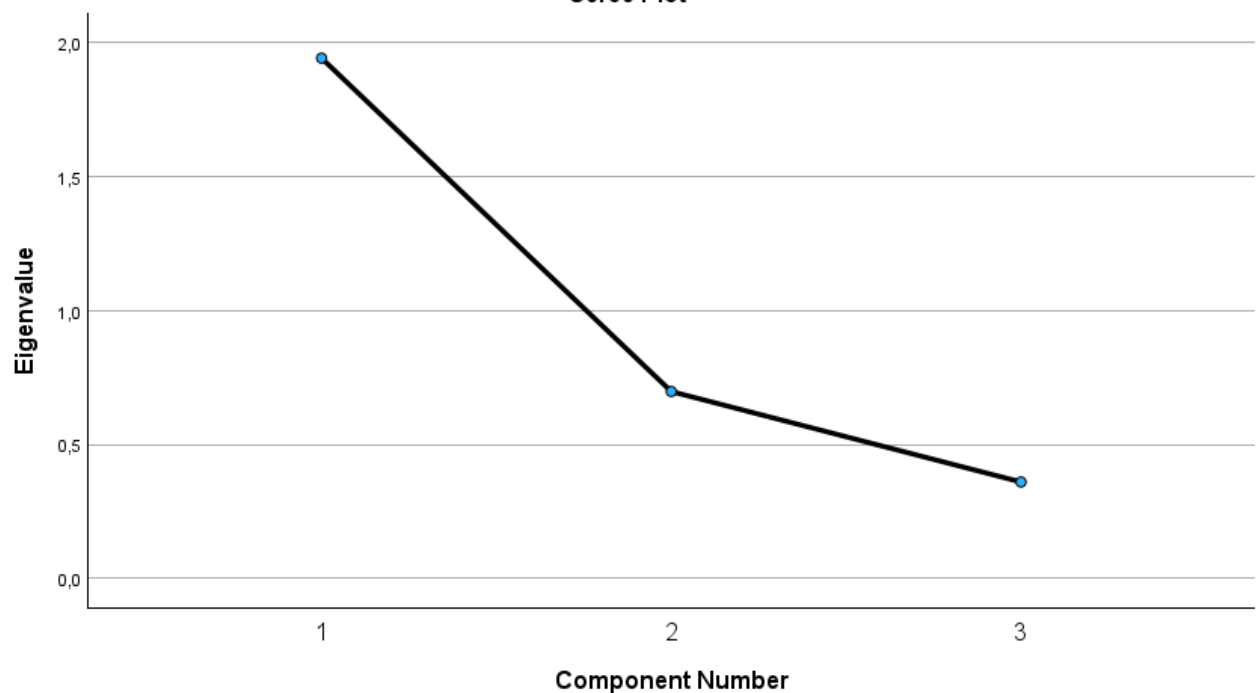
Extraction Method: Principal Component Analysis.

Total Variance Explained

Component	Total	Initial Eigenvalues		Extraction Sums of Squared Loadings		
		% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	1,943	64,751	64,751	1,943	64,751	64,751
2	,698	23,256	88,006			
3	,360	11,994	100,000			

Extraction Method: Principal Component Analysis.

Scree Plot



Component Matrix^a

	Component 1
[Manau, kad šis socialinių tinklų įrašas ir jo autorius atitinka vienas kitą.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,882
[Manau, kad šis socialinių tinklų įrašas yra susijęs su jo autoriumi.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,808
[Manau, kad socialinių tinklų įrašo autorius veiksmingai perteikė numatytą žinutę.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,716

Extraction Method: Principal Component Analysis.

a. 1 components extracted.

10 priedas. Faktorinės analizės rezultatai. Suvokiamas turinio patikimumas

Covariance Matrix

	[Šis socialinių tinklų įrašas kelia pasitikėjimą] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	[Šis socialinių tinklų įrašas atrodo sąžiningas] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	[Šis socialinių tinklų įrašas atrodo įtikinamas] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.
[Šis socialinių tinklų įrašas kelia pasitikėjimą] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	2,474	1,596	2,009
[Šis socialinių tinklų įrašas atrodo sąžiningas] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,596	2,034	1,678
[Šis socialinių tinklų įrašas atrodo įtikinamas] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	2,009	1,678	2,542

KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		,742
Bartlett's Test of Sphericity	Approx. Chi-Square	703,626
	df	3
	Sig.	<,001

Communalities

	Initial	Extraction
[Šis socialinių tinklų įrašas kelia pasitikėjimą] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,000	,843
[Šis socialinių tinklų įrašas atrodo sąžiningas] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,000	,795
[Šis socialinių tinklų įrašas atrodo įtikinamas] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	1,000	,862

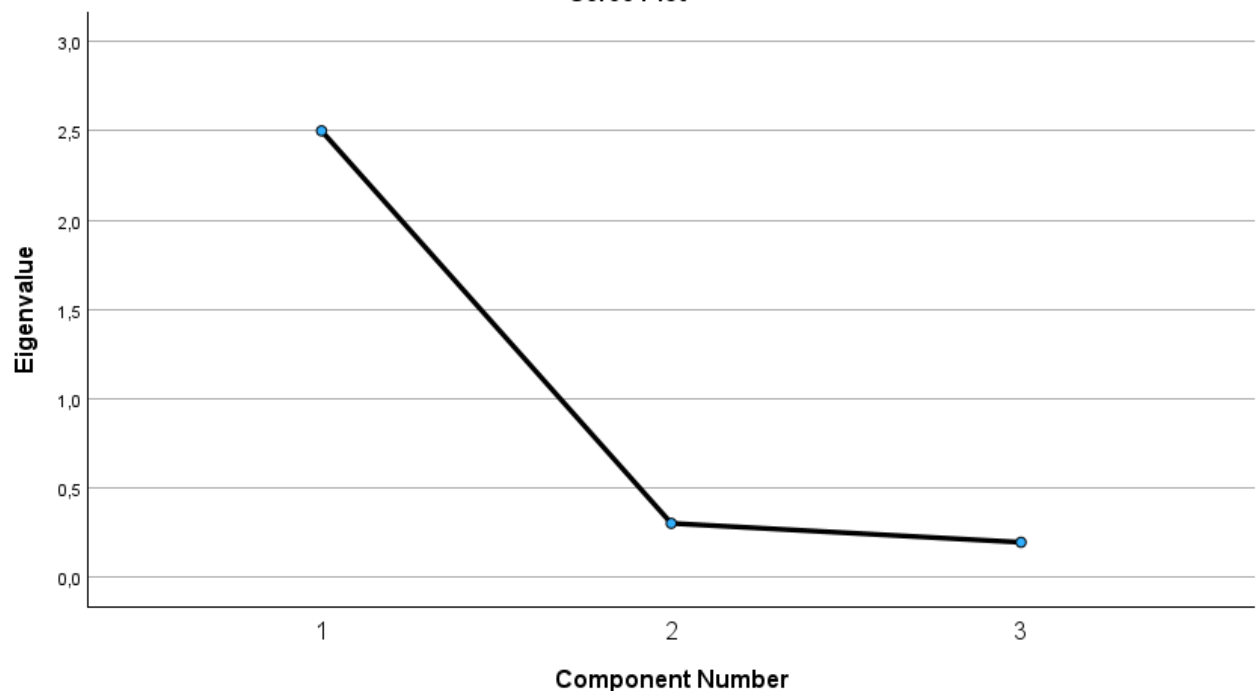
Extraction Method: Principal Component Analysis.

Total Variance Explained

Component	Total	Initial Eigenvalues		Extraction Sums of Squared Loadings		
		% of Variance	Cumulative %	Total	% of Variance	Cumulative %
1	2,501	83,365	83,365	2,501	83,365	83,365
2	,303	10,086	93,451			
3	,196	6,549	100,000			

Extraction Method: Principal Component Analysis.

Scree Plot



Component Matrix^a

	Component 1
[Šis socialinių tinklų įrašas atrodo įtikinamas] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,929
[Šis socialinių tinklų įrašas kelia pasitikėjimą] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,918
[Šis socialinių tinklų įrašas atrodo sąžiningas] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	,892

Extraction Method: Principal Component Analysis.

a. 1 components extracted.

11 priedas. Skalių patikimumo vertinimas. Šaltinio patikimumo per pasitikėjimą skalė

Case Processing Summary

		N	%
Cases	Valid	372	100,0
	Excluded ^a	0	,0
	Total	372	100,0

a. Listwise deletion based on all variables in the procedure.

Reliability Statistics

Cronbach's Alpha	N of Items
,903	3

Item Statistics

	Mean	Std. Deviation	N
[Sąžiningas Nesąžiningas] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau	4,65	1,458	372
[Garbingas Negarbingas] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau esate	4,55	1,508	372
[Moralus Amoralus] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau esate tikra	4,80	1,505	372

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
13,99	16,760	4,094	3

12 priedas. Skalių patikimumo vertinimas. Turinio patikimumo skalė

Case Processing Summary

		N	%
Cases	Valid	372	100,0
	Excluded ^a	0	,0
	Total	372	100,0

a. Listwise deletion based on all variables in the procedure.

Reliability Statistics

Cronbach's Alpha	N of Items
,900	3

Item Statistics

	Mean	Std. Deviation	N
[Šis socialinių tinklų įrašas kelia pasitikėjimą] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	3,98	1,573	372
[Šis socialinių tinklų įrašas atrodo sąžiningas] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	4,33	1,426	372
[Šis socialinių tinklų įrašas atrodo įtikinamas] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	4,24	1,594	372

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
12,55	17,617	4,197	3

13 priedas. Skalių patikimumo vertinimas. Šaltinio patikimumo per kompetenciją skalė

Case Processing Summary

	N	%
Cases		
Valid	372	100,0
Excluded ^a	0	,0
Total	372	100,0

a. Listwise deletion based on all variables in the procedure.

Reliability Statistics

Cronbach's Alpha	N of Items
,889	6

Item Statistics

	Mean	Std. Deviation	N
[Protingas Neprotingas] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau esate	4,85	1,486	372
[Neapmokytas Apmokytas] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau esate	4,83	1,581	372
[Ne ekspertas Ekspertas] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau esate	4,34	1,564	372
[Informuotas Neinformuotas] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau es	4,82	1,664	372
[Nekompetentingas Kompetentingas] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo lab	4,76	1,493	372
[Sumanus Kvailas] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau esate tikras	4,95	1,506	372

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
28,55	55,655	7,460	6

14 priedas. Skalių patikimumo vertinimas. Šaltinio patikimumo per palankumą skalė

Case Processing Summary

		N	%
Cases	Valid	372	100,0
	Excluded ^a	0	,0
	Total	372	100,0

a. Listwise deletion based on all variables in the procedure.

Reliability Statistics

Cronbach's Alpha	N of Items
,855	3

Item Statistics

	Mean	Std. Deviation	N
[Rūpinasi manimi Nesirūpina manimi] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo	4,04	1,600	372
[Rūpinasi mano interesais Nesirūpina mano interesais] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš	4,12	1,537	372
[Domisi manimi Nesidomi manimi] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labia	3,76	1,686	372

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
11,92	18,064	4,250	3

15 priedas. Skalių patikimumo vertinimas. Šaltinio patikimumo per empatiją skalė

Case Processing Summary

		N	%
Cases	Valid	372	100,0
	Excluded ^a	0	,0
	Total	372	100,0

a. Listwise deletion based on all variables in the procedure.

Reliability Statistics

Cronbach's Alpha	N of Items
,828	3

Item Statistics

	Mean	Std. Deviation	N
[Savanaudis Nesavanaudis] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau esat	4,62	1,522	372
[Nejautrus Jautrus] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau esate tikr	4,31	1,424	372
[Nesupratingas Supratingas] Išreikškite savo nuomonę apie socialinių tinklų įrašo autorių, pažymėdami savo pasirinkimą tarp pateiktų priešingų būdvardžių porų. Kuo arčiau jūsų pažymėjimas yra vieno iš būdvardžių, tuo labiau es	4,58	1,460	372

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
13,51	14,461	3,803	3

16 priedas. Skalių patikimumo vertinimas. Vartotojų įsitraukimo ketinimų skalė

Case Processing Summary

		N	%
Cases	Valid	372	100,0
	Excluded ^a	0	,0
	Total	372	100,0

a. Listwise deletion based on all variables in the procedure.

Reliability Statistics

Cronbach's Alpha	N of Items
,820	3

Item Statistics

	Mean	Std. Deviation	N
[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, paspausčiau „patinka“ po šiuo įrašu] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	3,34	1,777	372
[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, palikčiau komentarą po šiuo įrašu] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	2,03	1,337	372
[Perskaitęs(-usi) šį įrašą socialiniuose tinkluose, pasidalinčiau juo su kitais] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	2,53	1,582	372

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
7,90	16,437	4,054	3

17 priedas. Skalių patikimumo vertinimas. Šaltinio ir turinio atitikimo skalė

Case Processing Summary

		N	%
Cases	Valid	372	100,0
	Excluded ^a	0	,0
	Total	372	100,0

a. Listwise deletion based on all variables in the procedure.

Reliability Statistics

Cronbach's Alpha	N of Items
,722	3

Item Statistics

	Mean	Std. Deviation	N
[Manau, kad šis socialinių tinklų įrašas ir jo autorius atitinka vienas kitą.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	4,32	1,346	372
[Manau, kad socialinių tinklų įrašo autorius veiksmingai perteikė numatytą žinutę.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	4,67	1,345	372
[Manau, kad šis socialinių tinklų įrašas yra susijęs su jo autoriumi.] Nurodykite savo sutikimo ar nesutikimo lygį su kiekvienu žemiau pateikiamu teiginiu.	4,05	1,524	372

Scale Statistics

Mean	Variance	Std. Deviation	N of Items
13,03	11,449	3,384	3

18 priedas. Tyrimo kintamųjų aprašomoji analizė

Descriptive Statistics

	N	Minimum	Maximum	Mean	Std. Deviation
Šaltinio ir turinio atitikimas	372	1,33	7,00	4,3441	1,12788
Turinio patikimumas	372	1,00	7,00	4,1846	1,39909
Įsitraukimo ketinimai	372	1,00	7,00	2,6326	1,35142
Šaltinio kompetencija	372	1,00	7,00	4,7590	1,24337
Šaltinio palankumas	372	1,00	7,00	3,9749	1,41674
Šaltinio empatija	372	1,00	7,00	4,5036	1,26758
Šaltinio pasitikėjimas	372	1,00	7,00	4,6649	1,36464
Valid N (listwise)	372				

Tests of Normality

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Šaltinio ir turinio neatitikimas	,101	372	<,001	,974	372	<,001
Turinio patikimumas	,109	372	<,001	,965	372	<,001
Įsitraukimo ketinimai	,116	372	<,001	,926	372	<,001
Šaltinio kompetencija	,072	372	<,001	,970	372	<,001
Šaltinio palankumas	,152	372	<,001	,966	372	<,001
Šaltinio empatija	,125	372	<,001	,970	372	<,001
Šaltinio pasitikėjimas	,125	372	<,001	,956	372	<,001

a. Lilliefors Significance Correction

Correlations

			Šaltinio ir turinio atitikimas	Turinio patikimumas	Įsitraukimo ketinimai	Šaltinio kompetencija	Šaltinio palankumas	Šaltinio empatija	Šaltinio pasitikėjimas
Spearman's rho	Šaltinio ir turinio atitikimas	Correlation Coefficient	1,000	,457**	,244**	,405**	,292**	,436**	,364**
		Sig. (2-tailed)	.	<,001	<,001	<,001	<,001	<,001	<,001
		N	372	372	372	372	372	372	372
	Turinio patikimumas	Correlation Coefficient	,457**	1,000	,374**	,566**	,356**	,599**	,563**
		Sig. (2-tailed)	<,001	.	<,001	<,001	<,001	<,001	<,001
		N	372	372	372	372	372	372	372
	Įsitraukimo ketinimai	Correlation Coefficient	,244**	,374**	1,000	,283**	,400**	,455**	,221**
		Sig. (2-tailed)	<,001	<,001	.	<,001	<,001	<,001	<,001
		N	372	372	372	372	372	372	372
	Šaltinio kompetencija	Correlation Coefficient	,405**	,566**	,283**	1,000	,449**	,555**	,629**
		Sig. (2-tailed)	<,001	<,001	<,001	.	<,001	<,001	<,001
		N	372	372	372	372	372	372	372
	Šaltinio palankumas	Correlation Coefficient	,292**	,356**	,400**	,449**	1,000	,427**	,455**
		Sig. (2-tailed)	<,001	<,001	<,001	<,001	.	<,001	<,001
		N	372	372	372	372	372	372	372
	Šaltinio empatija	Correlation Coefficient	,436**	,599**	,455**	,555**	,427**	1,000	,570**
		Sig. (2-tailed)	<,001	<,001	<,001	<,001	<,001	.	<,001
		N	372	372	372	372	372	372	372
	Šaltinio pasitikėjimas	Correlation Coefficient	,364**	,563**	,221**	,629**	,455**	,570**	1,000
		Sig. (2-tailed)	<,001	<,001	<,001	<,001	<,001	<,001	.
		N	372	372	372	372	372	372	372

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

Ranks

	Jūs esate	N	Mean Rank
Šaltinio ir turinio atitikimas	Vyras	150	176,90
	Moteris	220	192,68
	Kita	2	226,50
	Total	372	
Turinio patikimumas	Vyras	150	173,13
	Moteris	220	196,86
	Kita	2	49,50
	Total	372	
Įsitraukimo ketinimai	Vyras	150	180,86
	Moteris	220	190,74
	Kita	2	143,50
	Total	372	
Šaltinio kompetencija	Vyras	150	175,49
	Moteris	220	194,85
	Kita	2	93,50
	Total	372	
Šaltinio palankumas	Vyras	150	171,45
	Moteris	220	196,85
	Kita	2	176,50
	Total	372	
Šaltinio empatija	Vyras	150	168,43
	Moteris	220	199,41
	Kita	2	121,50
	Total	372	
Šaltinio pasitikėjimas	Vyras	150	171,82
	Moteris	220	197,16
	Kita	2	114,50
	Total	372	

Test Statistics^{a,b}

	Šaltinio ir turinio atitikimas	Turinio patikimumas	Įsitraukimo ketinimai	Šaltinio kompetencija	Šaltinio palankumas	Šaltinio empatija	Šaltinio pasitikėjimas
Kruskal-Wallis H	2,222	7,662	1,087	4,411	5,078	8,248	5,955
df	2	2	2	2	2	2	2
Asymp. Sig.	,329	,022	,581	,110	,079	,016	,051

a. Kruskal Wallis Test

b. Grouping Variable: Jūs esate

Ranks

	Amžius_Kvart	N	Mean Rank
Šaltinio ir turinio atitikimas	15–23 m.	105	157,17
	24–25 m.	81	189,91
	26–31 m.	96	200,13
	32–67 m.	90	203,12
	Total	372	
Turinio patikimumas	15–23 m.	105	162,13
	24–25 m.	81	162,41
	26–31 m.	96	210,35
	32–67 m.	90	211,17
	Total	372	
Įsitraukimo ketinimai	15–23 m.	105	151,38
	24–25 m.	81	184,45
	26–31 m.	96	173,08
	32–67 m.	90	243,63
	Total	372	
Šaltinio kompetencija	15–23 m.	105	169,20
	24–25 m.	81	174,57
	26–31 m.	96	210,54
	32–67 m.	90	191,77
	Total	372	
Šaltinio palankumas	15–23 m.	105	164,86
	24–25 m.	81	197,39
	26–31 m.	96	182,48
	32–67 m.	90	206,23
	Total	372	
Šaltinio empatija	15–23 m.	105	163,13
	24–25 m.	81	186,13
	26–31 m.	96	184,50
	32–67 m.	90	216,23
	Total	372	
Šaltinio pasitikėjimas	15–23 m.	105	168,39
	24–25 m.	81	194,85
	26–31 m.	96	192,50
	32–67 m.	90	193,72
	Total	372	

Test Statistics^{a,b}

	Šaltinio ir turinio atitikimas	Turinio patikimumas	Įsitraukimo ketinimai	Šaltinio kompetencija	Šaltinio palankumas	Šaltinio empatija	Šaltinio pasitikėjimas
Kruskal-Wallis H	11,708	19,047	38,581	8,754	8,381	12,040	4,243
df	3	3	3	3	3	3	3
Asymp. Sig.	,008	<,001	<,001	,033	,039	,007	,236

a. Kruskal Wallis Test

b. Grouping Variable: Amžius_Kvart

Ranks

	Koks yra aukščiausias jūsų įgytas išsilavinimas?	N	Mean Rank
Šaltinio ir turinio atitikimas	Pradinis išsilavinimas	62	163,92
	Pagrindinis išsilavinimas	24	203,75
	Vidurinis išsilavinimas	212	188,07
	Aukštesnysis arba specialusis vidurinis išsilavinimas	6	198,83
	Aukštasis išsilavinimas	68	195,03
	Total	372	
Turinio patikimumas	Pradinis išsilavinimas	62	174,47
	Pagrindinis išsilavinimas	24	202,17
	Vidurinis išsilavinimas	212	187,94
	Aukštesnysis arba specialusis vidurinis išsilavinimas	6	222,50
	Aukštasis išsilavinimas	68	184,26
	Total	372	
Įsitraukimo ketinimai	Pradinis išsilavinimas	62	155,66
	Pagrindinis išsilavinimas	24	246,67
	Vidurinis išsilavinimas	212	183,54
	Aukštesnysis arba specialusis vidurinis išsilavinimas	6	286,17
	Aukštasis išsilavinimas	68	193,82
	Total	372	
Šaltinio kompetencija	Pradinis išsilavinimas	62	177,44
	Pagrindinis išsilavinimas	24	204,00
	Vidurinis išsilavinimas	212	187,41
	Aukštesnysis arba specialusis vidurinis išsilavinimas	6	250,83
	Aukštasis išsilavinimas	68	180,09
	Total	372	
Šaltinio palankumas	Pradinis išsilavinimas	62	162,50
	Pagrindinis išsilavinimas	24	255,00
	Vidurinis išsilavinimas	212	187,53
	Aukštesnysis arba specialusis vidurinis išsilavinimas	6	229,17
	Aukštasis išsilavinimas	68	177,24
	Total	372	
Šaltinio empatija	Pradinis išsilavinimas	62	171,02
	Pagrindinis išsilavinimas	24	233,42
	Vidurinis išsilavinimas	212	187,04
	Aukštesnysis arba specialusis vidurinis išsilavinimas	6	255,83
	Aukštasis išsilavinimas	68	176,26
	Total	372	
Šaltinio pasitikėjimas	Pradinis išsilavinimas	62	172,44
	Pagrindinis išsilavinimas	24	207,42
	Vidurinis išsilavinimas	212	190,02
	Aukštesnysis arba specialusis vidurinis išsilavinimas	6	241,83
	Aukštasis išsilavinimas	68	176,09
	Total	372	

Test Statistics ^{a,b}							
	Šaltinio ir turinio atitikimas	Turinio patikimumas	Įsitraukimo ketinimai	Šaltinio kompetencija	Šaltinio palankumas	Šaltinio empatija	Šaltinio pasitikėjimas
Kruskal-Wallis H	3,944	2,040	18,458	3,491	14,528	9,094	4,497
df	4	4	4	4	4	4	4
Asymp. Sig.	,414	,728	,001	,479	,006	,059	,343

a. Kruskal Wallis Test

b. Grouping Variable: Koks yra aukščiausias jūsų įgytas išsilavinimas?

Ranks

	Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?	N	Mean Rank
Šaltinio ir turinio atitikimas	Mažiau nei 500 Eur	26	173,35
	500–999 Eur	50	164,38
	1000–1499 Eur	114	184,38
	1500–1999 Eur	68	230,18
	2000–2499 Eur	26	193,96
	2500 Eur ar daugiau	30	155,63
	Nenoriu atsakyti	58	177,05
	Total	372	
Turinio patikimumas	Mažiau nei 500 Eur	26	204,50
	500–999 Eur	50	154,22
	1000–1499 Eur	114	192,64
	1500–1999 Eur	68	185,32
	2000–2499 Eur	26	227,88
	2500 Eur ar daugiau	30	174,23
	Nenoriu atsakyti	58	183,36
	Total	372	
Įsitraukimo ketinimai	Mažiau nei 500 Eur	26	161,04
	500–999 Eur	50	175,58
	1000–1499 Eur	114	188,31
	1500–1999 Eur	68	197,24
	2000–2499 Eur	26	219,88
	2500 Eur ar daugiau	30	186,37
	Nenoriu atsakyti	58	176,29
	Total	372	
Šaltinio kompetencija	Mažiau nei 500 Eur	26	222,12
	500–999 Eur	50	140,18
	1000–1499 Eur	114	203,92
	1500–1999 Eur	68	189,68
	2000–2499 Eur	26	202,65
	2500 Eur ar daugiau	30	145,77
	Nenoriu atsakyti	58	186,33
	Total	372	
Šaltinio palankumas	Mažiau nei 500 Eur	26	171,04
	500–999 Eur	50	176,22
	1000–1499 Eur	114	197,85
	1500–1999 Eur	68	192,82
	2000–2499 Eur	26	225,96
	2500 Eur ar daugiau	30	127,97
	Nenoriu atsakyti	58	185,16
	Total	372	
Šaltinio empatija	Mažiau nei 500 Eur	26	209,73
	500–999 Eur	50	169,02
	1000–1499 Eur	114	196,45
	1500–1999 Eur	68	180,68
	2000–2499 Eur	26	212,96
	2500 Eur ar daugiau	30	160,57
	Nenoriu atsakyti	58	179,98
	Total	372	
Šaltinio pasitikėjimas	Mažiau nei 500 Eur	26	163,27
	500–999 Eur	50	158,58
	1000–1499 Eur	114	203,48
	1500–1999 Eur	68	180,88
	2000–2499 Eur	26	226,65
	2500 Eur ar daugiau	30	157,97
	Nenoriu atsakyti	58	190,95
	Total	372	

Test Statistics ^{a,b}							
	Šaltinio ir turinio atitikimas	Turinio patikimumas	Įsitraukimo ketinimai	Šaltinio kompetencija	Šaltinio palankumas	Šaltinio empatija	Šaltinio pasitikėjimas
Kruskal-Wallis H	16,988	9,973	5,779	20,134	15,141	7,342	13,679
df	6	6	6	6	6	6	6
Asymp. Sig.	,009	,126	,448	,003	,019	,290	,033

a. Kruskal Wallis Test

b. Grouping Variable: Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?

19 priedas. Homogeniškumo testavimas

Jūs esate * {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)} Crosstabulation

			{if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}				Total
			Su žymėjimu, atitinka DI	Be žymėjimo, atitinka DI	Su žymėjimu, neatitinka DI	Be žymėjimo, neatitinka DI	
Jūs esate	Vyrai	Count	40	42	26	42	150
		% within Jūs esate	26,7%	28,0%	17,3%	28,0%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	41,7%	41,2%	40,6%	38,9%	40,5%
		% of Total	10,8%	11,4%	7,0%	11,4%	40,5%
	Moteris	Count	56	60	38	66	220
		% within Jūs esate	25,5%	27,3%	17,3%	30,0%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	58,3%	58,8%	59,4%	61,1%	59,5%
		% of Total	15,1%	16,2%	10,3%	17,8%	59,5%
Total	Count		96	102	64	108	370
	% within Jūs esate		25,9%	27,6%	17,3%	29,2%	100,0%
	% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}		100,0%	100,0%	100,0%	100,0%	100,0%
	% of Total		25,9%	27,6%	17,3%	29,2%	100,0%

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)
Pearson Chi-Square	,190 ^a	3	,979
Likelihood Ratio	,190	3	,979
Linear-by-Linear Association	,175	1	,676
N of Valid Cases	370		

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 25,95.

Symmetric Measures

		Value	Approximate Significance
Nominal by Nominal	Phi	,023	,979
	Cramer's V	,023	,979
N of Valid Cases		370	

Descriptives

Koks yra jūsų amžius?

	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
Su žymėjimu, atitinka DI	98	31,08	11,148	1,126	28,85	33,32	20	61
Be žymėjimo, atitinka DI	102	29,35	9,182	,909	27,55	31,16	18	61
Su žymėjimu, neatitinka DI	64	28,81	10,575	1,322	26,17	31,45	15	54
Be žymėjimo, neatitinka DI	108	28,40	9,492	,913	26,59	30,21	20	67
Total	372	29,44	10,076	,522	28,41	30,47	15	67

Tests of Homogeneity of Variances

		Levene Statistic	df1	df2	Sig.
Koks yra jūsų amžius?	Based on Mean	1,738	3	368	,159
	Based on Median	,764	3	368	,515
	Based on Median and with adjusted df	,764	3	354,255	,515
	Based on trimmed mean	1,523	3	368	,208

ANOVA

Koks yra jūsų amžius?

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	407,307	3	135,769	1,341	,261
Within Groups	37256,271	368	101,240		
Total	37663,578	371			

Robust Tests of Equality of Means

Koks yra jūsų amžius?

	Statistic ^a	df1	df2	Sig.
Welch	1,196	3	186,718	,313
Brown-Forsythe	1,317	3	321,212	,269

a. Asymptotically F distributed.

Multiple Comparisons

Dependent Variable: Koks yra jūsų amžius?

Tukey HSD

(I) (if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK))		(J) (if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK))	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
						Lower Bound	Upper Bound
Su žymėjimu, atitinka DI	Be žymėjimo, atitinka DI		1,729	1,423	,618	-1,94	5,40
	Su žymėjimu, neatitinka DI		2,269	1,617	,498	-1,90	6,44
	Be žymėjimo, neatitinka DI		2,683	1,404	,225	-,94	6,31
Be žymėjimo, atitinka DI	Su žymėjimu, atitinka DI		-1,729	1,423	,618	-5,40	1,94
	Su žymėjimu, neatitinka DI		,540	1,605	,987	-3,60	4,68
	Be žymėjimo, neatitinka DI		,955	1,389	,902	-2,63	4,54
Su žymėjimu, neatitinka DI	Su žymėjimu, atitinka DI		-2,269	1,617	,498	-6,44	1,90
	Be žymėjimo, atitinka DI		-,540	1,605	,987	-4,68	3,60
	Be žymėjimo, neatitinka DI		,414	1,587	,994	-3,68	4,51
Be žymėjimo, neatitinka DI	Su žymėjimu, atitinka DI		-2,683	1,404	,225	-6,31	,94
	Be žymėjimo, atitinka DI		-,955	1,389	,902	-4,54	2,63
	Su žymėjimu, neatitinka DI		-,414	1,587	,994	-4,51	3,68

Koks yra jūsų amžius?

Tukey HSD^{a,b}

{if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}		Subset for alpha = 0.05
	N	1
Be žymėjimo, neatitinka DI	108	28,40
Su žymėjimu, neatitinka DI	64	28,81
Be žymėjimo, atitinka DI	102	29,35
Su žymėjimu, atitinka DI	98	31,08
Sig.		,285

Means for groups in homogeneous subsets are displayed.

- Uses Harmonic Mean Sample Size = 89,102.
- The group sizes are unequal. The harmonic mean of the group sizes is used. Type I error levels are not guaranteed.

Ranks

{if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}		N	Vidurkių rangai
Koks yra aukščiausias jūsų įgytas išsilavinimas?	Su žymėjimu, atitinka DI	98	195,62
	Be žymėjimo, atitinka DI	102	183,23
	Su žymėjimu, neatitinka DI	64	178,34
	Be žymėjimo, neatitinka DI	108	186,15
	Total	372	

Test Statistics^{a,b}

Koks yra aukščiausias jūsų įgytas išsilavinimas?	
Kruskal-Wallis H	1,454
df	3
Asymp. Sig.	,693

- Kruskal Wallis Test
- Grouping Variable: {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}

Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)? * {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}
Crosstabulation

			{if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}				Total
			Su žymėjimu, atitinka DI	Be žymėjimo, atitinka DI	Su žymėjimu, neatitinka DI	Be žymėjimo, neatitinka DI	
Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?	Mažiau nei 500 Eur	Count	4	6	6	10	26
		% within Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?	15,4%	23,1%	23,1%	38,5%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	4,1%	5,9%	9,4%	9,3%	7,0%
		% of Total	1,1%	1,6%	1,6%	2,7%	7,0%
	500–999 Eur	Count	10	12	12	16	50
		% within Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?	20,0%	24,0%	24,0%	32,0%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	10,2%	11,8%	18,8%	14,8%	13,4%
		% of Total	2,7%	3,2%	3,2%	4,3%	13,4%
	1000–1499 Eur	Count	34	30	14	36	114
		% within Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?	29,8%	26,3%	12,3%	31,6%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	34,7%	29,4%	21,9%	33,3%	30,6%
		% of Total	9,1%	8,1%	3,8%	9,7%	30,6%
	1500–1999 Eur	Count	26	14	12	16	68
		% within Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?	38,2%	20,6%	17,6%	23,5%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	26,5%	13,7%	18,8%	14,8%	18,3%
		% of Total	7,0%	3,8%	3,2%	4,3%	18,3%
	2000–2499 Eur	Count	8	10	4	4	26
		% within Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?	30,8%	38,5%	15,4%	15,4%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	8,2%	9,8%	6,3%	3,7%	7,0%
		% of Total	2,2%	2,7%	1,1%	1,1%	7,0%
	2500 Eur ar daugiau	Count	4	10	8	8	30
		% within Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?	13,3%	33,3%	26,7%	26,7%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	4,1%	9,8%	12,5%	7,4%	8,1%
		% of Total	1,1%	2,7%	2,2%	2,2%	8,1%
	Nenoriu atsakyti	Count	12	20	8	18	58
		% within Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?	20,7%	34,5%	13,8%	31,0%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	12,2%	19,6%	12,5%	16,7%	15,6%
		% of Total	3,2%	5,4%	2,2%	4,8%	15,6%
Total		Count	98	102	64	108	372
		% within Kokios yra Jūsų mėnesinės pajamos (atskaičius mokesčius)?	26,3%	27,4%	17,2%	29,0%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	100,0%	100,0%	100,0%	100,0%	100,0%
		% of Total	26,3%	27,4%	17,2%	29,0%	100,0%

Case Processing Summary

	Valid		Cases Missing		Total	
	N	Percent	N	Percent	N	Percent
Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose? * {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	372	100,0%	0	0,0%	372	100,0%

Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose? * {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)} Crosstabulation

		{if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}				Total
		Su žymėjimu, atitinka DI	Be žymėjimo, atitinka DI	Su žymėjimu, neatitinka DI	Be žymėjimo, neatitinka DI	
Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose?	1 valandą ir mažiau	Count	26	16	4	54
		% within Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose?	48,1%	29,6%	7,4%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	26,5%	15,7%	6,3%	14,5%
	2 valandas	Count	28	28	24	118
		% within Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose?	23,7%	23,7%	20,3%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	28,6%	27,5%	37,5%	31,7%
	3 valandas	Count	18	42	26	114
		% within Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose?	15,8%	36,8%	22,8%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	18,4%	41,2%	40,6%	30,6%
	4 valandas	Count	20	4	4	52
		% within Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose?	38,5%	7,7%	7,7%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	20,4%	3,9%	6,3%	14,0%
	5 valandas ir daugiau	Count	6	12	6	34
		% within Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose?	17,6%	35,3%	17,6%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	6,1%	11,8%	9,4%	9,1%
Total		Count	98	102	64	372
		% within Kiek vidutiniškai laiko per dieną praleidžiate socialiniuose tinkluose?	26,3%	27,4%	17,2%	100,0%
		% within {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	100,0%	100,0%	100,0%	100,0%

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)
Pearson Chi-Square	49,980 ^a	12	<,001
Likelihood Ratio	52,570	12	<,001
Linear-by-Linear Association	6,093	1	,014
N of Valid Cases	372		

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 5,85.

20 priedas. Manipuliacijos analizė

Pažymėta, kad įrašą sukūrė DI * Įrašas buvo pažymėtas dirbtinio intelekto autorystę nurodančia etikete „AI info“ (liet. dirbtinio intelekto informacija). Crosstabulation

			Įrašas buvo pažymėtas dirbtinio intelekto autorystę nurodančia etikete „AI info“ (liet. dirbtinio intelekto informacija).			
			Taip	Ne	Nežinau	Total
Pažymėta, kad įrašą sukūrė DI	Ne	Count	0	94	116	210
		% within Pažymėta, kad įrašą sukūrė DI	0,0%	44,8%	55,2%	100,0%
		% within Įrašas buvo pažymėtas dirbtinio intelekto autorystę nurodančia etikete „AI info“ (liet. dirbtinio intelekto informacija).	0,0%	100,0%	54,2%	56,5%
		% of Total	0,0%	25,3%	31,2%	56,5%
	Taip	Count	64	0	98	162
		% within Pažymėta, kad įrašą sukūrė DI	39,5%	0,0%	60,5%	100,0%
		% within Įrašas buvo pažymėtas dirbtinio intelekto autorystę nurodančia etikete „AI info“ (liet. dirbtinio intelekto informacija).	100,0%	0,0%	45,8%	43,5%
		% of Total	17,2%	0,0%	26,3%	43,5%
	Total	Count	64	94	214	372
		% within Pažymėta, kad įrašą sukūrė DI	17,2%	25,3%	57,5%	100,0%
		% within Įrašas buvo pažymėtas dirbtinio intelekto autorystę nurodančia etikete „AI info“ (liet. dirbtinio intelekto informacija).	100,0%	100,0%	100,0%	100,0%
		% of Total	17,2%	25,3%	57,5%	100,0%

Report

Vertinimas

randgroup	Mean	N	Std. Deviation
Su žymėjimu, atitinka DI	4,86	98	1,450
Be žymėjimo, atitinka DI	4,31	102	1,573
Su žymėjimu, neatitinka DI	4,94	64	1,531
Be žymėjimo, neatitinka DI	4,87	108	1,773
Total	4,73	372	1,610

ANOVA

Vertinimas

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	24,136	3	8,045	3,157	,025
Within Groups	937,896	368	2,549		
Total	962,032	371			

ANOVA Effect Sizes^{a,b}

			95% Confidence Interval	
Point Estimate			Lower	Upper
Vertinimas	Eta-squared	,025	,000	,058
	Epsilon-squared	,017	-,008	,050
	Omega-squared Fixed-effect	,017	-,008	,050
	Omega-squared Random-effect	,006	-,003	,017

a. Eta-squared and Epsilon-squared are estimated based on the fixed-effect model.

b. Negative but less biased estimates are retained, not rounded to zero.

Multiple Comparisons

Dependent Variable: Vertinimas

Tukey HSD

(I) {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	(J) {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Su žymėjimu, atitinka DI	Be žymėjimo, atitinka DI	,543	,226	,078	-,04	1,13
	Su žymėjimu, neatitinka DI	-,080	,257	,989	-,74	,58
	Be žymėjimo, neatitinka DI	-,013	,223	1,000	-,59	,56
Be žymėjimo, atitinka DI	Su žymėjimu, atitinka DI	-,543	,226	,078	-1,13	,04
	Su žymėjimu, neatitinka DI	-,624	,255	,070	-1,28	,03
	Be žymėjimo, neatitinka DI	-,557	,220	,058	-1,13	,01
Su žymėjimu, neatitinka DI	Su žymėjimu, atitinka DI	,080	,257	,989	-,58	,74
	Be žymėjimo, atitinka DI	,624	,255	,070	-,03	1,28
	Be žymėjimo, neatitinka DI	,067	,252	,993	-,58	,72
Be žymėjimo, neatitinka DI	Su žymėjimu, atitinka DI	,013	,223	1,000	-,56	,59
	Be žymėjimo, atitinka DI	,557	,220	,058	-,01	1,13
	Su žymėjimu, neatitinka DI	-,067	,252	,993	-,72	,58

Vertinimas

Tukey HSD^{a,b}

{if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}		Subset for alpha = 0.05	
	N	1	2
Be žymėjimo, atitinka DI	102	4,31	
Su žymėjimu, atitinka DI	98	4,86	4,86
Be žymėjimo, neatitinka DI	108	4,87	4,87
Su žymėjimu, neatitinka DI	64		4,94
Sig.		,094	,987

Means for groups in homogeneous subsets are displayed.

- a. Uses Harmonic Mean Sample Size = 89,102.
 b. The group sizes are unequal. The harmonic mean of the group sizes is used. Type I error levels are not guaranteed.

Ranks

{if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}		N	Mean Rank
Vertinimas	Su žymėjimu, atitinka DI	98	193,13
	Be žymėjimo, atitinka DI	102	157,60
	Su žymėjimu, neatitinka DI	64	200,34
	Be žymėjimo, neatitinka DI	108	199,57
	Total	372	

Test Statistics^{a,b}

Vertinimas	
Kruskal-Wallis H	10,789
df	3
Asymp. Sig.	,013

- a. Kruskal Wallis Test
 b. Grouping Variable: {if(is_empty(randgroup.NAOK), rand(1,4), randgroup.NAOK)}

Atvirų respondentų atsakymų pasiskirstymas pagal scenarijų

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Su žymėjimu, atitinka DI	48	26,7	26,7	26,7
	Be žymėjimo, atitinka DI	48	26,7	26,7	53,3
	Su žymėjimu, neatitinka DI	30	16,7	16,7	70,0
	Be žymėjimo, neatitinka DI	54	30,0	30,0	100,0
	Total	180	100,0	100,0	

21 priedas. Atvirų atsakymų analizė

Nr.	Autorystės vertinimas	Emocinis tonas	Atsakymas
1	Neaišku	Neutralus	Šiuolaikinis dirbtinis intelektas tikrai gali parašyti tokį tekstą, tačiau lygiai taip pat gali ir žmogus
2	Abu	Neutralus	Žmogus kartu su DI, tai bendras darbas
3	DI	Neutralus	Emoji atrodo išdėlioti ne taip kaip žmogus dėtų
4	Abu	Neigiamas	Tekstas galbūt ir parašytas žmogaus, tačiau labai informacinis ir nuobodus, o nuotrauka tikrai sugeneruota DI
5	Neaišku	Teigiamas	Sunku atskirti ar sukurta žmogaus, ar dirbtinio intelekto. Tekstas sklandus, aiškiai suprantamas
6	DI	Neutralus	Viršuje parašyta "AI info"
7	DI	Neutralus	Nuotrauka atrodo sugeneruota dirbtinio intelekto, todėl iškart atrodo, kad tekstinis įrašas nėra rašytas žmogaus.
8	DI	Neutralus	Sunku pasakyti, bet senjorės atvaizdas atrodo sukurtas AI, jeigu ne AI atrodo kad panaudota labai daug photoshop.
9	DI	Neigiamas	Nuotrauka atrodo labai dirbtiniai
10	Abu	Neutralus	Vaizdas galbūt ir sukurtas dirbtinio intelekto, bet tikrai ne tekstas.
11	DI	Neutralus	Matosi, kad paveikslėlis tikrai sukurtas DI
12	Žmogus	Teigiamas	Prie įrašo nėra nurodyta, kad jį sukūrė dirbtinis intelektas. Matydama, kad įrašas yra iš mokslo įstaigos – laikau įrašą patikimesniu ir manau, kad komunikuoja universiteto atstovai.
13	Neaišku	Teigiamas	Tekstas atrodo paprastas, nesudėtingas be įmantrių terminų. Lengvai skaitosi.
14	DI	Neigiamas	Atrodo kažkoks nenatūralus, matosi, kad nuotrauka tikrai DI
15	Neaišku	Neigiamas	suardomas rankomis - kas čia per pasakymas :D
16	Abu	Neutralus	Mano patirtyje DI arba naudoja per daug "emoji", arba jų išvis nenaudoja. DI geriausiai rašo anglų kalba, tekstą generuotą lietuvių kalba reiktų modifikuoti, ištaisyti klaidas. Todėl, nors šis įrašas gali būti sukurtas DI, prie jo 100% prisidėjo žmogus.
17	Žmogus	Neutralus	Manau jog parašė žmogus.
18	DI	Neutralus	Bendra teksto suformuluota struktūra labai primena dirbtinio intelekto rašymo būdą.
19	DI	Neutralus	Prie įrašo pažymėta "AI info"
20	DI	Neutralus	Paveikslėlis atrodo taip, lyg būtų sugeneruotas su DI

21	Neaišku	Neutralus	Sunku suprasti
22	DI	Neigiamas	Tikėtina, nes yra nelabai žmogaus kalbai būdingų išsireiškimų (pavyzdžiui “veido įrašų”), tekstas nėra logiškas ir iki galo suprantamas.
23	Neaišku	Neutralus	Tekstas labai techniškas, sunku pasakyti, nes jis rodos, neatspindi jokios emocijos, tik perduoda informaciją.
24	DI	Neutralus	Pateikta nuotrauka atrodo, kaip DI darbas. Pats įrašas nėra labai akivaizdus, tačiau tie emoji, prieš įrašo pradžią, kelia įtarimų
25	DI	Neigiamas	Mano nuomone, žmogus nesirinktų į tokio tipo įrašą įtraukti parinktų emociukų. Be to, tekstas neskamba labai rišliai, labiau kaip faktų susakymas.
26	DI	Neutralus	Iš odos struktūros
27	DI	Neutralus	Nes aiškiai matosi
28	DI	Neutralus	Po KTU parašyta Ai info bei nuotrauka atrodo nereali. Nežinau ar KTU naudotų ? emoji, spėju, kad ne :D
29	DI	Neigiamas	Laužyta ir nerišli kalba
30	Abu	Neigiamas	50/50. Pradžia skamba natūraliau, nes sakiniai ilgesnis, tikroviškesni (lyg būtų rašyti žmogaus). Paskutiniai du sakiniai labiau skamba kaip DI darbas, nes ne taip "gražiai" (sklandžiai) sujungti su likusia teksto dalimi - tiesiog pasako faktą trumpai ir aiškiai (beeeeet trumpai!, o tas šiek tiek ir sukelia abejonių). O nuotrauka tai taip - sugeneruota DI (nors visai atrodo tikroviškai, bet still something is off :))
31	DI	Neutralus	Man taip pat kraujas tirštas, bet su vandens gėrimu nesusidraugauju
32	Neaišku	Neutralus	Dirbtinis intelektas gali padėti idėją tobulinti. Bet pati idėja manau yra žmogaus
33	Neaišku	Teigiamas	atrodo patikimai
34	Abu	Neutralus	Manau, kad pats įrašo pagrindas galėjo būti sukurtas DI, bet žmogus redagavo tekstą. Gal pridėjo emoji, gal sutvarkė skyrybą.
35	DI	Neutralus	Panašu į DI tekstą
36	DI	Neutralus	Sunku teigti vienaip ar kitaip, įrašas galėtų būti sukurtas dirbtinio intelekto
37	Neaišku	Neutralus	Nuotrauka atrodo, lyg galėtų būt padaryta AI, bet iš KTU tikėčiausi AI label'io, tai sunku pasakyti:D
38	DI	Neigiamas	Vizualas 100% sukurtas dirbtinio intelekto programėlėmis, tą išduoda tiksliai vizuale išdėstytos situacijos ir situacijų veikėjai, vis dar negebėjimas atvaizduoti žmogaus galūnes(pirštus). Dėl teksto turinio komentuoti negaliu.
39	Neaišku	Neutralus	Nelabai supratau su kuo čia reikia sutikti ar nesutikti
40	Abu	Teigiamas	Nuotrauka netikroviška bet tekstas normalus
41	Žmogus	Teigiamas	Sunku pasakyti, atrodo žmogaus rašytas, ne “plastikinis” tekstas
42	DI	Neutralus	Lakoniškas tekstas, panašų tekstą tikrai galima gauti paprašius AI parašyti apie naują kuriamą produktą (nors galėtų būti nebūtinai AI kurtas). Nuotrauka atrodo galėtų būti kurta AI pagalba

43	Žmogus	Teigiamas	Per daug realistiška
44	Abu	Neutralus	Gali būti ir žmogus, ir AI, visi tie marketinginiai tekstai panašūs:)
45	DI	Neutralus	Viršuje indikacija AI info
46	DI	Neigiamas	Nuotrauka tai labai "perspausta"
47	DI	Neigiamas	DI generuojamas nuotraukas dažnai paslysta ties žmonių pirštais, tad čia ne išimtis.
48	Abu	Neutralus	Dėl teksto sunku pasakyti, bet nuotrauka tikrai sukurta su dirbtiniu intelektu, tai nurodo priedas "AI info".
49	Neaišku	Neutralus	Sutinku, nes tai problema
50	Neaišku	Neigiamas	Nepasitikiu tokiais tyrimais, nes daug suklustama..
51	Neaišku	Neutralus	Abejonių kelia foto, bet tiksliai nežinau
52	DI	Neigiamas	Pateikiama daug faktų, tekstas gana struktūruotas. Iš esmės nuobodžiai skaitosi toks vientisas teksto blokas
53	DI	Neutralus	AI info
54	DI	Neigiamas	Šabloniniai sakiniai
55	Neaišku	Neigiamas	Atrodo, kaip nerealus produktas.
56	DI	Neigiamas	Paveikslėlio detalės kai kurios atrodo truputį netikros, o parašytas tekstas šiek tiek šabloninis, panašus į tokį, kokį pats gaunu naudodamas DI
57	Abu	Neutralus	Manau parašė žmogus pagrindines mintis, bet redakcijai galėjo būti panaudotas DI
58	DI	Teigiamas	Atrodo taip tobulai, kad net kyla įtarimas, kad netikra
59	Abu	Neutralus	Sunku įvertinti. Klaidų nesimato, tad matosi, jog papildoma redakcija tikrai atlikta, tačiau gali būti, kad griaučius parašė DI.
60	Abu	Neutralus	Atrodo, kad tekstas parašytas DI ir pakoreguotas žmogaus.
61	Neaišku	Teigiamas	Aprašyta be klaidų, logiškai
62	DI	Neutralus	Nurodyta, jog tai "AI info"
63	Neaišku	Teigiamas	Geras produktas, ne per daug siūlymo
64	Neaišku	Neutralus	Nežinau ar tai tiesa
65	DI	Neutralus	Daug turinio soc tinkluose sukurtas DI
66	DI	Neigiamas	Išreikšta mokslinė kalba esantis tekstas bei panaudoti papildomi priedai tested, kas straipsnių rašymuose papildomi filtrai (paveikslėliai) nėra naudojami reklamos pateikimui. Tai neprofesionalu.

67	Neaišku	Neutralus	Mokslo pažanga akivaizdi naudojant DI.
68	Abu	Teigiamas	Tekstas atrodo konkretus apie aktualią medžiagą ir problemą, tačiau paveikslėlis galimai yra sukurtas DI
69	Žmogus	Teigiamas	Gana sklandu
70	Neaišku	Neutralus	nesu susipažinusi su šia problema
71	DI	Neutralus	Kai DI programos kaip ChatGPT kuria tekstą, jos labai dažnai panaudoja ilgąjį brūkšnį vadinamu em-dash. Nors tai nėra neabejotinas DI naudojimo įrodymas, tai taip pat yra labai pastebima DI savybė.
72	DI	Teigiamas	Labai sklandžiai aprašyta situacija, manau tai DI aprašas.
73	Neaišku	Neutralus	Produktas pagamintas iš 95% vandens tačiau yra tvirtas, bet tuo pačiu ir lengvai tirpsta. Šios savybės skamba prieštarškai tarpusavio atžvilgiu. Iškyla klausimas kaip pasisavinamas burnoje ištirpęs gaminys, be ryjimo.
74	DI	Neigiamas	keista sakinių struktūra, tekstas lyg verstas iš anglų kalbos
75	Žmogus	Neigiamas	Nes kalbama apie konkretų universiteto pasiekimą, apie kurį DI nežinotų, kartu dažnai su soc tinkluose DI įrašais poruojamos DI sugeneruoti vaizdai ir šiame įraše pateiktas vaizdas panašus į nelabai moksliškai surengtą fotosesiją
76	Neaišku	Neutralus	Pirmam sakinyje paminėtas dirbtinio intelekto sprendimas
77	Neaišku	Neutralus	Mokslininkai sukūrė naują DI sprendimą. Tik remtasi DI, pats DI be suteikiamų komandų negeba kurti tokių dalykų, tuo labiau fizinių.
78	Neaišku	Neigiamas / teigiamas	Paveikslas nelogiškas, tačiau tekstas rišlus.
79	DI	Neigiamas	Ai info, rašybos klaidos, emoji.
80	DI	Neutralus	Neretai DI pateikus užduotį sugeneruoti tekstą socialiniams tinklams, jis pradedamas „jaustuku“ (angl. emoji) ir kartojasi svarbiausiose (labiausiai pabrėžtinose) tekso vietose, kad atkreipti skaitančiųjų dėmesį.
81	Abu	Neutralus	Įrašas tikriausiai kurtas žmogaus, o nuotrauka - dirbtinio intelekto
82	DI	Neigiamas	Keistos struktūros pastraipa (keisti minčių šuoliai joje). Apskritai skamba kaip paturbintas marketinginio slang'o vėmalas arba bent jau kaip jį bandytų pamėgdžioti populiarūs AI modeliai. +Atrodo kaip AI daromas vertimas iš anglų kalbos (per daug literaliai verčiama, nėra "lokalizacijos" jausmo).
83	Žmogus	Neutralus	Man šis tekstas skambėjo kaip sukurtas žmogaus.
84	Abu	Neutralus	Nuotrauka labai panaši į DI generuotą, tačiau negaliu to teigti 100 procentų. Įrašo aprašymas panašus, kad rašytas žmogaus.
85	Neaišku	Neutralus	Atrodo kažkaip neįtikima ir per paprasta, kad depresiją galima būtų nustatyti pasitelkiant balsą.

86	DI	Neutralus	Daug vardinamų aspektų per kablelį, emoji.
87	Abu	Neutralus	Nuotrauka matosi, kad sugeneruota, o formalų tekstą sunku atskirti žmogaus ar AI kurtas.
88	Neaišku	Neutralus	vandens trukumas ir rijimo problema, vandeni kaip tik lengva nuryti, dar buckio emoji,
89	Abu	Neutralus	Nuotrauka DI, nuojauda sako, kad tekstas galėjo būti readaguojamas su DI
90	Neaišku	Neutralus	Sunku atskirti.
91	DI	Neigiamas	Nenatūraliai sudėlioti sakiniai ir neįprasti žodžiai.
92	Neaišku	Neutralus	Atrodo paprastas tekstas.
93	DI	Neigiamas	Labai daug terminų klaidų, aiškiai AI sugeneruota nuotrauka, trumpi sakiniai su labai daug apibūdinimų, bet tuo pačiu tekstas sausas.
94	Neaišku	Neutralus	Nesu tikra ar ji sukūrė dirbtinis intelektas.
95	DI	Neigiamas	akivaizdžiai matosi išlietas vaizdas
96	DI	Neigiamas	Trūksta atributikos minėtams mokslininkams, jų projektui. Nėra nuorodų, specifikų apie taikomas chemijos technologijas. Emotikonai atrodo kaip blogai "fintuned" modelio neprisitaikymas prie konteksto (universiteto pranešimo), kaip pilkos socialinės medijos paskyros sensacija. Projekto idėja absurdiška -- bandoma sukurti vandens pakaitalą (solidų) kuris vis tiek tirpsta burnoje, nors mediciniškai springimas skysčiais yra gan retas ir tikriausiai nėra aktualus žmonėms kurie iš viso dar gali rankas judinti.
97	Žmogus	Neutralus	Nuotrauka atrodo natūrali nederbtinai sugeneruota. Tekstas išsamus, be bendračių, be punktų iš didžiųjų raidžių.
98	Žmogus	Neutralus	Nes remiantis posto autoriumi susidarau nuomonė, kad tai vieno iš studentu parašytas tekstas.
99	DI	Neutralus	Tekste naudojami tam tikri žodžiais ir perėjimai nuo minties prie minties, kurie sudaro įspūdį jog tai sukūrė dirbtinis intelektas
100	DI	Neutralus	Foto generuota DI
101	DI	Neigiamas	Viena ranka turi 4 pirštus, tai pat dirbtinio intelekto kurti vaizdai turi "švelnius bruožus"
102	Žmogus	Teigiamas	Tekstas sklandus, aiškus, išsamus. Nenaudojami hiperbolizmai ar nerodoma per daug entuziazmo, kuriuo pasižymi DI.
103	DI	Neigiamas	Žmogui AI programai pateikė padrikas mintis, todėl pats tekstas nesudėliotas rišliai
104	Abu	Neutralus	Manau, kad DI tekstuose nededa emoji. Tačiau nuotrauka kelia įtarimų, jo ji būtent yra sukurta DI įrankiu.
105	DI	Neigiamas	Emoji naudojimas, ilgesni brūkšniai būdingi DI, nuotrauka atrodo netikroviškai.
106	Neaišku	Neutralus	Sunku pasakyti
107	DI	Neigiamas	Nes mikrofoną laikantis žmogus žiūri niekur, o jo kairės rankos pirštų - daugiau nei turėtų būti
108	DI	Neutralus	Sakinių struktūra primena DI sugeneruotą tekstą.

109	DI	Neigiamas	Fone monitoriuje bereikšmis vaizdas
110	Žmogus	Neutralus	Dėl to, kad teksto autorius yra institucija - KTU, plius jau taip aprašomas mokslinis DL tyrimas.
111	Neaišku	Neutralus	Neturiu argumento
112	Neaišku	Neutralus	Taip atrodo. Galbūt universitetas nesirinktų, kad DI kurtų jiems tekstus
113	Neaišku	Neutralus	negaliu atsakyti į šį klausimą.
114	DI	Neigiamas	Tekstas yra labai mechaninis, nevisiškai rišlus. Atrodo lyg būtų vertinys iš anglų kalbos.
115	DI	Neigiamas	Lietuvių kalba rašomas tekstas kai kuriose vietose skamba nenatūraliai.
116	DI	Neigiamas	Nuotrauka atrodo netikra, įrašo aprašas atrodo labai lakoniškai, neatrodo, kad paremtas straipsniais ir t.t.
117	Abu	Teigiamas	Tekstas parašytas tvarkingai, aiškiai. Tą gali tiek DI, tiek žmogus.
118	DI	Neigiamas	Nuotraukos dizainas, tekstas tarsi "negyvas"
119	DI	Neutralus	Tekstas parasytas tokia sakinių struktūra, kurią naudoja dirbtinis intelektas.
120	DI	Neigiamas	Neįtikina, kad produktas gali būti sudarytas iš 95% vandens, nes atrodo kaip guminukas. Be to, nuotrauka neatrodo autentiška
121	DI	Neigiamas	Pirmiausiai dėl emoji naudojimo. Antra, iš pirmos pažiūros tekstas skamba gražiai, tačiau dekonstruojant pastraipą, jau pirmajame sakinyje galima pastebėti nelogišką priežasties-pasekmės sąsają: ar tikrai KTU sukūrė įrenginį, nes depresija paliečia milijonus žmonių. Turbūt priežastis yra, nes norima jiems pagelbėti.
122	Žmogus	Teigiamas	Parašyta rišliai ir aiškiai suprantamai.
123	Abu	Neutralus	Gali būti ir taip ir taip
124	DI	Neigiamas	Atrodo kaip AI sugeneruota, nes spalvos nėra tokios ryškios
125	DI	Neigiamas	Asmuo antrame fone atrodo netikras
126	DI	Neigiamas	kvaila nuotrauka
127	DI	Neigiamas	Odos tekstūra per daug lygi ir glotni, plaukai atrodo gana tvarkingi, visa aplinka per daug estetiška ir lygi, tvarkingai. Fonas yra "blure" arba nesufokusuotas, pastebėjau, kad tai dažnas bruožas DI sukurto vaizdo. Pamačius nuotrauką ir vėliau skaitant tekstą, buvo kilusi mintis, kad nuotrauka sugeneravo dirbtinis intelektas, tai tikliausiai ir tekstą jis rašė.

128	Abu	Neutralus	Manau, kad turint informaciją apie šį sukūrimą, apibendrinimą šiam įrašui galėjo parašyti asmuo. Bet tobulėjant dirbtiniam intelektui ir jam tinkamai suformulavus užduotį bei pateikus visą informaciją yra šansas, kad įrašas galėjo būti padarytas su DI.
129	DI	Neutralus	Nuotrauka daug ką pasako
130	Žmogus	Neutralus	Nepanasu , kad cia kurtu AI
131	DI	Neutralus	AI info etiketė
132	Neaišku	Neutralus	Nesu tikra
133	DI	Neigiamas	Nuotrauka DI nes yra blur efektas fone, saldainiai taip pat atrodo per daug tobuli ir lygūs kaip ir plaukai bei veidas. Saldainio laikymas atrodo nenatūralus.
134	DI	Neigiamas	Nuotrauka sugeneruota DI (vienos ranko keturi pirštai). Tekste esantys emoji (pvz, pirmasis emoji ir pats pirmas, nors dažniausiai žmogui parašius jis eitu po sakinio, ne prieš jį.
135	DI	Neigiamas	Pirštai neorganiški, susiliejimas su kede truputi ir rubu
136	DI	Neutralus	Nuotrauka tikrai sudaryta AI
137	DI	Neutralus	Mano nuomone, naudojami „emoji“ tekste sukelia įspūdį, kad tekstas buvo generuojamas pasitelkiant DI.
138	Abu	Neutralus	Įrašo vizualą panašu, kad kūrė DI, tekstą, galbūt, ir žmogus.
139	Neaišku	Neutralus	Sukūrė intelektas bet diagnostika komentuoja specialistas
140	Neaišku	Teigiamas	Jeį dirbtinis intelektas, skamba gan naturaliai, nedirbtinai, faktai pateikiami žmogiškai
141	Žmogus	Neutralus	Neatrodo
142	Abu	Neutralus	Tikiu kad šį tekstą galėjo sukurti dirbtinis intelektas, bet neatmenu ir galimybės kad tai buvo parašyta žmogaus be dirbtinio intelekto pagalbos.
143	DI	Neigiamas	skamba nerealistiškai, nes atrodo, kad jei užkandžio didžioji dalis sudaryta iš vandens, tai kaip tik turėtų būti lengva suardyti rankomis?
144	DI	Neutralus	Nuotrauka atrodo sukurta dirbtinio intelekto
145	DI	Neutralus	Emoji ir brūkšnių naudojimas bei pats išdėstymas tekste kelia įtarimą, nes taip daro dažnai dirbtinis intelektas
146	Abu	Neigiamas	Sunku pasakyti.. yra gramatikos klaidų,tai manyčiau,kad tiesiog žmogiškasis faktorius.. bet DI irgi turbūt galėtų daryti gramatinių klaidų ?
147	DI	Neigiamas	Nuotrauka atrodo dirbtina
148	DI	Neutralus	Paveikslėlis vizualiai atrodo sugeneruotas AI, pradžioje teksto panaudotas emoji, taip kaip tą daro AI programos.

149	DI	Neutralus	Pagal nuojautą
150	Abu	Neutralus	Tekstas atrodo parašytas žmogiškuoju stiliumi, tačiau nuotrauka, panašu, jog būtų generuota dirbtinio intelekto pagalba.
151	DI	Neigiamas	Emoji keistai integruoti į sakinius.
152	DI	Neigiamas	Žmogaus plaukai susipina su laidais,
153	Žmogus	Teigiamas	Atrodo labai natūraliai.
154	Žmogus	Teigiamas	Tekstas parašytas, atsižvelgiant į respondento įtraukimą, naudojamos socialiai svarios frazės ("Lietuvoje analogų neturintį"), žinant, kad tai pakreips nuomonę į teigiamą pusę. Idėti vaizdūs apibūdinimai, o ne techninės specifikacijos. Tekstas sudėliotas taip, kad sukeltų teigiamas reakcijas ir vaizduotėje produkto prototipą, nėra vien techniškas ir faktiškas.
155	DI	Neigiamas	Tekstas neskamba "žmogiškai", atrodo toks... perdėtai išstobulintas, nuotrauka irgi neįtikina
156	DI	Neutralus	Vizualas sukurtas dirbtinio intelekto, matosi iš senjorės rankų
157	DI	Neutralus	Nuotrauka ryškiai sugeneruota ir tekstas nėra tipinio ktu formato
158	Neaišku	Neutralus	Nežinau ką atsakyti
159	DI	Neigiamas	Eigu vertinti nuotrauką: smegenų veiklos davikliai uždėti vienam žmogui, bet mikrofonas pas kita žmogų. Iš teksto suprantu, kad tyrime turėtų vertinti vieno žmogaus tiek smegenų veiklą tiek žmogaus kalbą, balso intonaciją. Yra nesutapimų tarp teksto ir nuotraukos.
160	Žmogus	Neutralus	Nematau nieko dirbtinio?
161	Neaišku	Neutralus	Tai galėjo parašyti ir žmogus
162	DI	Neigiamas	Gramatikos, Stiliaus klaidos
163	Neaišku	Neigiamas	Sudetinga būtų patikėti DI savo psichine sveikata. Labiau pasitikejimo klausimas. Baisu, kad nepadarytu blogiau.
164	DI	Neigiamas	Atrodo dirbtina
165	Žmogus	Teigiamas	Sklandi kalba,
166	DI	Neutralus	Iš pačio teksto yra sunku nuspresti ar tai sukūrė dirbtinis intelektas, tačiau naudojama nuotrauka primena butent AI braiza (spalvos, žmonių veidai)
167	DI	Neigiamas	Yra specifin? sakinio struktūra, būdinga AI ir kai kurie žodžiai nevisai tinkamai panaudoti

168	Žmogus	Teigiamas	Pateiktas oficialus KTU soc.tinklo pavadinimas (atrodo, kad tai verified account); nėra gramatinių klaidų, tekstas įtikinamas, panašus į KTU įrašų stilių (nors šiek tiek ilgokas, bet nesusimąstyčiau, kad šį tekstą kūrė DI); nuotrauka kiek keista, atrodo netikroviškai, bet tiesiog vartydama FB nepagalvočiau, kad tai DI.
169	DI	Neigiamas	Žodžiai atrodo neorganiškai sukurti, o ir taip žmonės geriau sugeba sukurti tokius panašius tekstus, jog pritrauktų skaitytoja
170	DI	Neigiamas	Sakiniai nėra labai rišliai sudėlioti, naudojamos keistos sąvokos. Nuotrauka nelabai atrodo realistiška
171	Neaišku	Neutralus	Nežinau, o kaip atskirti, kad įrašą sukūrė DI?
172	DI	Neigiamas	Kai kurios formuluotės tokios kaip "o balsas pasirinktas kaip" kiek keistokos
173	Abu	Neutralus	Neaišku ar apklausoje kalbama apie vaizdą ar paį tekstą. Nes vaizdas galimai kurtas di, o tekstas ne di
174	DI	Neigiamas	Dešinės rankos pirštai, atrodo, kad jų tik keturi, tuomet iš paveikslo kokybės matosi
175	Žmogus	Neutralus	Nepažymėta AI
176	DI	Neigiamas	Nuotrauka atrodo labai redaguota ir pats tekstas neskamba tikroviškai ar lyg būtų rašytas žmogaus. Kasnio dydžio užkandis niekaip neatitiks 2-3 reikalingų litrų vandens per dieną ir apskritai negali būti siūlomas žmogui, kuris turi rijimo sutrikimų. Taip pat, minimi "svarbūs" mineralai kaip selenis, geležis ir cinkas, bet visiškai neminimi reikalingi elektrolitai - natrij, kalis ir magnis.
177	DI	Neutralus	Nors naudojami emoji's, tekstas daugmaž atrodo lyg būtų parašytas dirbtinio intelekto, nes naudojama labiau dalykinė kalba, neperteikiama daug emocijų. Drįsčiau tvirtinti, kad nuotrauka tikriausiai sukurta dirbtinio intelekto, nes žmonės atrodo dirbtinai. Beje ant kairiojo asmens laktos pritvirtinto elektrodo nėra prijungtų laidų, o plaukai kiaurai praeina pro elektrodą. Atrodo nelogiška.
178	DI	Neigiamas	Nes kalbama apie dar nerealų produktą, nuotrauka atrodo nenatūrali
179	DI	Neigiamas	Neaišku apie ką kalbama, iš pradžių apie vandenį, po to apie rijimo sunkumus, vėliau apie vitaminus. Neaiški žinutė, dėk ko kilo abejonių ar žmogus ją parašė
180	Neaišku	Neutralus	Kauno technologijos universiteto mokslininkai sukūrė naują DI sprendimą.

22 priedas. Hipotezių testavimas

H1 hipotezės testavimas.

Group Statistics

	Pažymėta, kad įrašą sukūrė DI	N	Mean	Std. Deviation	Std. Error Mean
Šaltinio ir turinio atitikimas	Ne	210	4,1619	1,12306	,07750
	Taip	162	4,5802	1,09302	,08588

Independent Samples Test

		Levene's Test for Equality of Variances		t-test for Equality of Means						95% Confidence Interval of the Difference	
		F	Sig.	t	df	One-Sided p	Two-Sided p	Mean Difference	Std. Error Difference	Lower	Upper
Šaltinio ir turinio atitikimas	Equal variances assumed	,002	,963	-3,604	370	<,001	<,001	-,41834	,11608	-,64660	-,19008
	Equal variances not assumed			-3,617	350,794	<,001	<,001	-,41834	,11568	-,64585	-,19084

Independent Samples Effect Sizes

		Standardizer ^a	Point Estimate	95% Confidence Interval	
				Lower	Upper
Šaltinio ir turinio atitikimas	Cohen's d	1,11009	-,377	-,583	-,170
	Hedges' correction	1,11235	-,376	-,582	-,170
	Glass's delta	1,09302	-,383	-,591	-,173

a. The denominator used in estimating the effect sizes.

Cohen's d uses the pooled standard deviation.

Hedges' correction uses the pooled standard deviation, plus a correction factor.

Glass's delta uses the sample standard deviation of the control (i.e., the second) group.

H2 hipotezės testavimas.

Descriptive Statistics

Dependent Variable: Šaltinio ir turinio atitikimas

Pažymėta, kad įrašą sukūrė DI	Sąsaja su DI	Mean	Std. Deviation	N
Ne	Nėra	4,1173	1,05342	108
	Yra	4,2092	1,19581	102
	Total	4,1619	1,12306	210
Taip	Nėra	4,2188	1,20730	64
	Yra	4,8163	,94538	98
	Total	4,5802	1,09302	162
Total	Nėra	4,1550	1,11076	172
	Yra	4,5067	1,11982	200
	Total	4,3441	1,12788	372

Levene's Test of Equality of Error Variances^{a,b}

		Levene Statistic	df1	df2	Sig.
Šaltinio ir turinio atitikimas	Based on Mean	2,328	3	368	,074
	Based on Median	2,007	3	368	,113
	Based on Median and with adjusted df	2,007	3	350,182	,113
	Based on trimmed mean	2,358	3	368	,071

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Dependent variable: Šaltinio ir turinio atitikimas

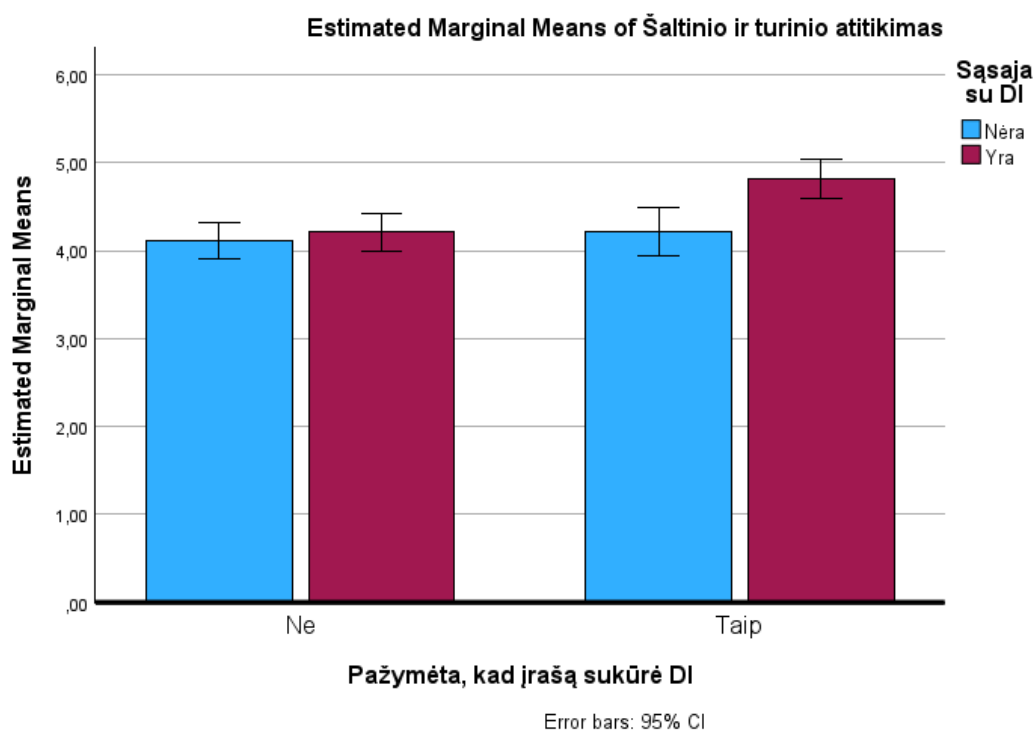
b. Design: Intercept + Autorystėsžyma + SąsajasuDI + Autorystėsžyma * SąsajasuDI

Tests of Between-Subjects Effects

Dependent Variable: Šaltinio ir turinio atitikimas

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
Corrected Model	30,273 ^a	3	10,091	8,408	<,001	,064
Intercept	6714,343	1	6714,343	5594,223	<,001	,938
Autorystėsžyma	11,186	1	11,186	9,320	,002	,025
SąsajasuDI	10,588	1	10,588	8,822	,003	,023
Autorystėsžyma * SąsajasuDI	5,697	1	5,697	4,746	,030	,013
Error	441,684	368	1,200			
Total	7492,000	372				
Corrected Total	471,957	371				

a. R Squared = ,064 (Adjusted R Squared = ,057)



H3 hipotezės testavimas. Kompetencijos dimensija

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	R Square Change	Change Statistics			
						F Change	df1	df2	Sig. F Change
1	,445 ^a	,198	,196	1,11472	,198	91,575	1	370	<,001

a. Predictors: (Constant), Šaltinio ir turinio atitikimas

b. Dependent Variable: Šaltinio kompetencija

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	113,791	1	113,791	91,575	<,001 ^b
	Residual	459,763	370	1,243		
	Total	573,553	371			

a. Dependent Variable: Šaltinio kompetencija

b. Predictors: (Constant), Šaltinio ir turinio atitikimas

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients Beta	t	Sig.	95,0% Confidence Interval for B		Correlations		
		B	Std. Error				Lower Bound	Upper Bound	Zero-order	Partial	Part
1	(Constant)	2,626	,230		11,403	<,001	2,173	3,079			
	Šaltinio ir turinio atitikimas	,491	,051	,445	9,569	<,001	,390	,592	,445	,445	,445

a. Dependent Variable: Šaltinio kompetencija

H3 hipotezės testavimas. Palankumo dimensija**Model Summary^b**

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	R Square Change	Change Statistics			
						F Change	df1	df2	Sig. F Change
1	,315 ^a	,099	,097	1,34659	,099	40,662	1	370	<,001

a. Predictors: (Constant), Šaltinio ir turinio atitikimas

b. Dependent Variable: Šaltinio palankumas

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	73,733	1	73,733	40,662	<,001 ^b
	Residual	670,921	370	1,813		
	Total	744,655	371			

a. Dependent Variable: Šaltinio palankumas

b. Predictors: (Constant), Šaltinio ir turinio atitikimas

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients Beta	t	Sig.	95,0% Confidence Interval for B		Correlations		
		B	Std. Error				Lower Bound	Upper Bound	Zero-order	Partial	Part
1	(Constant)	2,258	,278		8,117	<,001	1,711	2,805			
	Šaltinio ir turinio atitikimas	,395	,062	,315	6,377	<,001	,273	,517	,315	,315	,315

a. Dependent Variable: Šaltinio palankumas

H3 hipotezės testavimas. Empatijos dimensija

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	R Square Change	Change Statistics			
						F Change	df1	df2	Sig. F Change
1	,461 ^a	,212	,210	1,12648	,212	99,761	1	370	<,001

a. Predictors: (Constant), Šaltinio ir turinio atitikimas

b. Dependent Variable: Šaltinio empatija

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	126,592	1	126,592	99,761	<,001 ^b
	Residual	469,514	370	1,269		
	Total	596,106	371			

a. Dependent Variable: Šaltinio empatija

b. Predictors: (Constant), Šaltinio ir turinio atitikimas

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients Beta	t	Sig.	95,0% Confidence Interval for B		Correlations		
		B	Std. Error				Lower Bound	Upper Bound	Zero-order	Partial	Part
1	(Constant)	2,254	,233		9,685	<,001	1,796	2,711			
	Šaltinio ir turinio atitikimas	,518	,052	,461	9,988	<,001	,416	,620	,461	,461	,461

a. Dependent Variable: Šaltinio empatija

H3 hipotezės testavimas. Pasitikėjimo dimensija**Model Summary^b**

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	R Square Change	Change Statistics			
						F Change	df1	df2	Sig. F Change
1	,373 ^a	,139	,137	1,26795	,139	59,735	1	370	<,001

a. Predictors: (Constant), Šaltinio ir turinio atitikimas

b. Dependent Variable: Šaltinio pasitikėjimas

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	96,036	1	96,036	59,735	<,001 ^b
	Residual	594,852	370	1,608		
	Total	690,888	371			

a. Dependent Variable: Šaltinio pasitikėjimas

b. Predictors: (Constant), Šaltinio ir turinio atitikimas

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients Beta	t	Sig.	95,0% Confidence Interval for B		Correlations		
		B	Std. Error				Lower Bound	Upper Bound	Zero-order	Partial	Part
1	(Constant)	2,705	,262		10,328	<,001	2,190	3,220			
	Šaltinio ir turinio atitikimas	,451	,058	,373	7,729	<,001	,336	,566	,373	,373	,373

a. Dependent Variable: Šaltinio pasitikėjimas

H4 hipotezės testavimas

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	R Square Change	Change Statistics			
						F Change	df1	df2	Sig. F Change
1	,450 ^a	,202	,200	1,25134	,202	93,783	1	370	<,001

a. Predictors: (Constant), Šaltinio ir turinio atitikimas

b. Dependent Variable: Turinio patikimumas

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	146,850	1	146,850	93,783	<,001 ^b
	Residual	579,364	370	1,566		
	Total	726,214	371			

a. Dependent Variable: Turinio patikimumas

b. Predictors: (Constant), Šaltinio ir turinio atitikimas

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients Beta	t	Sig.	95,0% Confidence Interval for B		Correlations		
		B	Std. Error				Lower Bound	Upper Bound	Zero-order	Partial	Part
1	(Constant)	1,761	,258		6,814	<,001	1,253	2,270			
	Šaltinio ir turinio atitikimas	,558	,058	,450	9,684	<,001	,445	,671	,450	,450	,450

a. Dependent Variable: Turinio patikimumas

H5 hipotezės testavimas**Model Summary^b**

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	R Square Change	Change Statistics			
						F Change	df1	df2	Sig. F Change
1	,216 ^a	,047	,044	1,32119	,047	18,173	1	370	<,001

a. Predictors: (Constant), Šaltinio ir turinio atitikimas

b. Dependent Variable: Įsitraukimo ketinimai

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	31,722	1	31,722	18,173	<,001 ^b
	Residual	645,846	370	1,746		
	Total	677,569	371			

a. Dependent Variable: Įsitraukimo ketinimai

b. Predictors: (Constant), Šaltinio ir turinio atitikimas

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients Beta	t	Sig.	95,0% Confidence Interval for B		Correlations		
		B	Std. Error				Lower Bound	Upper Bound	Zero-order	Partial	Part
1	(Constant)	1,506	,273		5,519	<,001	,970	2,043			
	Šaltinio ir turinio atitikimas	,259	,061	,216	4,263	<,001	,140	,379	,216	,216	,216

a. Dependent Variable: Įsitraukimo ketinimai

H6 ir H7 hipotezių testavimas

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate	R Square Change	Change Statistics			
						F Change	df1	df2	Sig. F Change
1	,536 ^a	,287	,277	1,14899	,287	29,448	5	366	<,001

a. Predictors: (Constant), Šaltinio pasitikėjimas, Šaltinio palankumas, Turinio patikimumas, Šaltinio empatija, Šaltinio kompetencija

b. Dependent Variable: Įsitraukimo ketinimai

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	194,383	5	38,877	29,448	<,001 ^b
	Residual	483,185	366	1,320		
	Total	677,569	371			

a. Dependent Variable: Įsitraukimo ketinimai

b. Predictors: (Constant), Šaltinio pasitikėjimas, Šaltinio palankumas, Turinio patikimumas, Šaltinio empatija, Šaltinio kompetencija

Coefficients^a

Model		Unstandardized Coefficients		Standardized Coefficients	t	Sig.	95,0% Confidence Interval for B		Correlations			Collinearity Statistics	
		B	Std. Error	Beta			Lower Bound	Upper Bound	Zero-order	Partial	Part	Tolerance	VIF
1	(Constant)	,525	,261		2,015	,045	,013	1,038					
	Turinio patikimumas	,292	,059	,303	4,940	<,001	,176	,409	,389	,250	,218	,519	1,927
	Šaltinio kompetencija	-,086	,072	-,079	-1,189	,235	-,227	,056	,240	-,062	-,052	,445	2,249
	Šaltinio palankumas	,241	,051	,253	4,747	<,001	,141	,341	,335	,241	,210	,686	1,459
	Šaltinio empatija	,359	,066	,337	5,402	<,001	,228	,489	,435	,272	,238	,502	1,992
	Šaltinio pasitikėjimas	-,275	,064	-,278	-4,323	<,001	-,401	-,150	,156	-,220	-,191	,471	2,123

a. Dependent Variable: Įsitraukimo ketinimai

H8 ir H9 hipotezių testavimas

Run MATRIX procedure:

Copyright 2013-2024 by Andrew F. Hayes. ALL RIGHTS RESERVED.

This version of PROCESS requires SPSS version 26 or later

Workshop schedule available at haskayne.ucalgary.ca/CCRAM

In SPSS 29 and later, change default output font to Courier New for tidier output. More information about PROCESS at processmacro.org/faq.html.

***** PROCESS Procedure for SPSS Version 5.0 beta 2.1 *****

Written by Andrew F. Hayes, Ph.D. www.afhayes.com

Documentation available in Hayes (2022). www.guilford.com/p/hayes3

Model: 4

Y: Isitrau

X: Atitik

M1: Tpatik

M2: Skompet

M3: Spalank

M4: Sempati

M5: Spasitik

Sample
Size: 372

OUTCOME VARIABLE:
Tpatik

Model Summary

R	R-sq	MSE	F	df1	df2	p
,4497	,2022	1,5658	93,7827	1,0000	370,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	1,7614	,2585	6,8141	,0000	1,2531	2,2697
Atitik	,5578	,0576	9,6841	,0000	,4445	,6711

OUTCOME VARIABLE:
Skompet

Model Summary

R	R-sq	MSE	F	df1	df2	p
,4454	,1984	1,2426	91,5745	1,0000	370,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	2,6259	,2303	11,4035	,0000	2,1731	3,0787
Atitik	,4910	,0513	9,5695	,0000	,3901	,5919

OUTCOME VARIABLE:
Spalank

Model Summary

R	R-sq	MSE	F	df1	df2	p
,3147	,0990	1,8133	40,6625	1,0000	370,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	2,2579	,2782	8,1169	,0000	1,7109	2,8049
Atitik	,3953	,0620	6,3767	,0000	,2734	,5171

OUTCOME VARIABLE:

Sempati

Model Summary

R	R-sq	MSE	F	df1	df2	p
,4608	,2124	1,2690	99,7608	1,0000	370,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	2,2537	,2327	9,6851	,0000	1,7962	2,7113
Atitik	,5179	,0519	9,9880	,0000	,4159	,6199

OUTCOME VARIABLE:

Spasitik

Model Summary

R	R-sq	MSE	F	df1	df2	p
,3728	,1390	1,6077	59,7346	1,0000	370,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	2,7053	,2619	10,3284	,0000	2,1902	3,2203
Atitik	,4511	,0584	7,7288	,0000	,3363	,5659

OUTCOME VARIABLE:

Isitrau

Model Summary

R	R-sq	MSE	F	df1	df2	p
,5359	,2872	1,3232	24,5143	6,0000	365,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	,5751	,2866	2,0068	,0455	,0115	1,1387
Atitik	-,0263	,0625	-,4206	,6743	-,1492	,0966
Tpatik	,2967	,0601	4,9348	,0000	,1785	,4149
Skompet	-,0814	,0727	-1,1202	,2633	-,2244	,0615
Spalank	,2429	,0510	4,7598	,0000	,1426	,3433
Sempati	,3636	,0675	5,3892	,0000	,2309	,4963
Spasitik	-,2755	,0638	-4,3209	,0000	-,4009	-,1501

***** TOTAL EFFECT MODEL *****

OUTCOME VARIABLE:

Isitrau

Model Summary

R	R-sq	MSE	F	df1	df2	p
,2164	,0468	1,7455	18,1735	1,0000	370,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	1,5064	,2729	5,5194	,0000	,9697	2,0431
Atitik	,2593	,0608	4,2630	,0000	,1397	,3788

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y

Effect	se	t	p	LLCI	ULCI
,2593	,0608	4,2630	,0000	,1397	,3788

Direct effect of X on Y

Effect	se	t	p	LLCI	ULCI
-,0263	,0625	-,4206	,6743	-,1492	,0966

Indirect effect(s) of X on Y:

	Effect	BootSE	BootLLCI	BootULCI
TOTAL	,2855	,0475	,1950	,3828
Tpatik	,1655	,0420	,0920	,2565
Skompet	-,0400	,0380	-,1233	,0270
Spalank	,0960	,0267	,0495	,1541
Sempati	,1883	,0469	,1050	,2882
Spasitik	-,1243	,0386	-,2136	-,0586

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:

95,0000

Number of bootstrap samples for bias-corrected bootstrap confidence intervals:

5000

----- END MATRIX -----