



Kauno technologijos universitetas

Ekonomikos ir verslo fakultetas

**Vartotojų polinkio priskirti moralinį veiksnumą DI ir
žmogaus sprendimams poveikis pasitenkinimui paslaugomis**

Baigiamasis magistro projektas

Miglė Povilaitytė

Projekto autorė

Doc. dr. Beata Šeinauskienė

Vadovė

Kaunas, 2025



Kauno technologijos universitetas

Ekonomikos ir verslo fakultetas

Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikis pasitenkinimui paslaugomis

Baigiamasis magistro projektas

Marketingo valdymas (6211LX038)

Miglė Povilaitytė

Projekto autorė

Doc. dr. Beata Šeinauskienė

Vadovė

Doc. dr. Rimantė Hopenienė

Recenzentė

Kaunas, 2025



Kauno technologijos universitetas

Ekonomikos ir verslo fakultetas

Miglė Povilaitytė

Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikis pasitenkinimui paslaugomis

Akademinio sąžiningumo deklaracija

Patvirtinu, kad:

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

Miglė Povilaitytė

Patvirtinta elektroniniu būdu

Povilaitytė, Miglė. Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikis pasitenkinimui paslaugomis. Magistro baigiamasis projektas / vadovė doc. dr. Beata Šeinauskienė; Kauno technologijos universitetas, Ekonomikos ir verslo fakultetas.

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

Reikšminiai žodžiai: moralinis veiksnumas, vartotojų pasitenkinimas, dirbtinis intelektas, paslaugos, sprendimo palankumas, suvokiamas sąžiningumas.

Kaunas, 2025. 104 p.

Santrauka

Baigiamuoju magistro projektu siekta išsiaiškinti vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikį pasitenkinimui paslaugomis. Globalios rinkos tendencijos atskleidžia, jog įvairiuose paslaugų sektoriuose, vis sparčiau diegiamos DI sistemos, siekiant padidinti veiklos našumą (BCG, 2024). Verslai pabrėžia, jog vienas svarbiausių rodiklių, vertinant DI strategijos efektyvumą, yra vartotojų patirties ir pasitenkinimo gerinimas (IBM, 2024). Vis dėlto, esminiu barjeru išlieka su DI taikymu susiję etiniai iššūkiai (IBM, 2024). Didėjanti šia technologija paremtų sprendimų plėtra, tose verslo srityse, kur vartotojai tiesiogiai sąveikauja su DI sistemomis, kelia didžiausią verslų susirūpinimą. Nėra žinoma, kaip kintanti vartotojų ir paslaugą teikiančių žmonių sąveika veikia vartotojų atsaką į DI sprendimus, ypač moraliai svarbiose paslaugų teikimo aplinkybėse. Be to, trūksta žinių apie vartotojų atsaką, kai šiuos sprendimus priima žmogus ir DI. Nors vartotojų sąveikos su DI sistemomis poveikio supratimas yra svarbus, vis dar mažai tyrinėta, kaip vertinami identiški, moraliai abejotini žmogaus ir DI sprendimai. Ankstesni moksliniai darbai pateikia įrodymų, jog žmogaus ir DI sąveika moraliai svarbiuose kontekstuose turi įtakos pasitikėjimui, tačiau pasitikėjimas, lyginant su pasitenkinimu, turi mažesnės įtakos teigiamiems vartotojų ketinimams ir įmonių finansiniams rezultatams (Söderlund, 2023a). Esami tyrimai nepateikia apibendrintų išvadų, kam, moraliai dviprasmiškose aplinkybėse, vartotojai bus linkę priskirti didesnę moralinį veiksnumą (DI ar žmogui) ir kaip šis polinkis veikia pasitenkinimą paslaugomis. Siekiant atsakyti į šiuos probleminius klausimus, apibrėžtas tyrimo objektas, tikslas, uždaviniai ir metodai. Tyrimo objektas – vartotojų priskiriamo moralinio veiksnumo DI ir žmogaus sprendimams poveikis pasitenkinimui paslaugomis. Tyrimo tikslas – teoriškai ir empiriškai pagrįsti vartotojų polinkį priskirti moralinį veiksnumą DI ir žmogaus sprendimams ir jo poveikį pasitenkinimui paslaugomis. Tyrimo uždaviniai: 1. Atskleisti vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis aktualumą ir problematiką; 2. Teoriškai argumentuoti conceptualų modelį, paaiškinantį vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikį pasitenkinimui paslaugomis; 3. Sudaryti vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis, tyrimo metodologiją; 4. Empiriškai patikrinti vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikį pasitenkinimui paslaugomis; 5. Pateikti vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikį pasitenkinimui paslaugomis, praktines rekomendacijas ir tolimesnių tyrimų kryptis. Kiekybinis, eksperimento dizaino tyrimas atskleidė, jog moraliniu požiūriu abejotinosiose paslaugų aplinkybėse, vartotojų polinkis priskirti moralumą nėra veikiamas sprendimo priėmėjo tipo – DI ar žmogaus. DI ir žmogaus gebėjimas elgtis, laikantis moralės normų, vertinamas panašiai. Polinkis priskirti veiksnumą reikšmingai didesnis žmogaus sprendimo atveju. Polinkį priskirti moralinį veiksnumą reikšmingai

teigiamai veikia sprendimo palankumas. Esant palankiam sprendimui, moralinis veiksnumas padidėja, nepalankiam – priešingai, tačiau šis ryšys nepriklauso nuo to, ar sprendimą priėmė žmogus, ar DI. Pastebėta, jog suvokiamas DI veiksnumas palankaus ir nepalankaus sprendimo atveju išlieka panašus, mažesnis nei žmogaus. Polinkis priskirti moralinį veiksnumą turi reikšmingą teigiamą poveikį suvokiamam sprendimo priėmėjo sąžiningumui. Vis dėlto, sprendimo priėmėjo tipas šio ryšio nemoderuoja. Įrodyta, kad suvokiamas sprendimo priėmėjo sąžiningumas teigiamai veikia pasitenkinimą paslauga, nepaisant sprendimo priėmėjo, DI ar žmogaus. Suvokiamas sąžiningumas tarpininkauja ryšiui tarp polinkio priskirti moralumą ir pasitenkinimo paslauga. Tačiau, suvokiamas sąžiningumas nemedijuoja ryšio tarp polinkio priskirti veiksnumą ir pasitenkinimo. Nors hipotetizuota nebuvo, pastebėtas tiesioginis, teigiamas ryšys tarp polinkio priskirti moralumą ir vartotojų pasitenkinimo, veiksnio atveju, šis ryšys yra taip pat reikšmingas, tačiau neigiamos krypties.

Povilaitytė, Miglė. The Impact of Consumers' Propensity to Attribute Moral Agency Towards AI and Human Decisions on Service Satisfaction. Master's Final Degree Project / supervisor Assoc. Prof. Beata Šeinauskienė; School of Economics and Business, Kaunas University of Technology.

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

Keywords: moral agency, consumer satisfaction, artificial intelligence, services, decision favourability, perceived fairness.

Kaunas, 2025. 104 p.

Summary

The final Master's project aimed to investigate the impact of consumers' propensity to attribute moral agency towards AI and human decisions on service satisfaction. Global market trends indicate that AI systems are being increasingly deployed in various service sectors to improve operational efficiency (BCG, 2024). Businesses emphasise that one of the most important indicators for measuring the effectiveness of an AI strategy is improving customer experience and satisfaction (IBM, 2024). Nevertheless, a major barrier remains the ethical challenges associated with the implementation of AI (IBM, 2024). Business concerns are most acute with the increasing development of solutions based on this technology in areas of business where users directly interact with AI systems. It is unknown how the changing interaction patterns between consumers and service providers affect consumer reactions to AI solutions, especially in morally relevant service contexts. Furthermore, limited knowledge is gained about how consumers react when humans and AI make these decisions. While understanding the impact of consumer interactions with AI systems is important, there is still little research on how identical, morally questionable Human or AI decisions are assessed. Previous research provides evidence that human-AI interactions in morally relevant contexts affect trust; however, compared to satisfaction, trust has a smaller impact on consumers' positive intentions and firms' financial performance (Söderlund, 2023a). Existing research does not provide generalised conclusions on who, in morally ambiguous circumstances, consumers will tend to attribute higher moral agency (AI or human) and how this propensity affects satisfaction with services. To answer these problematic questions, the study's object, aim, objectives, and methods were defined. The object of the study is the impact of the moral agency attributed by consumers towards AI and human decisions on satisfaction with services. The study aims to provide theoretical and empirical justification for consumers' propensity to attribute moral agency towards AI and human decisions and its impact on service satisfaction. Objectives of the study: 1. To uncover the relevance and problematics of the impact of consumers' propensity to attribute moral agency towards AI and human decisions on service satisfaction; 2. To theoretically justify a conceptual model explaining the impact of consumers' propensity to attribute moral agency towards AI and human decisions on service satisfaction; 3. To develop a research methodology to investigate the impact of consumers' propensity to attribute moral agency towards AI and human decisions on service satisfaction; 4. To empirically test the impact of consumers' propensity to attribute moral agency towards AI and human decisions on service satisfaction; 5. To provide practical recommendations and directions for further research on the impact of consumers' propensity to attribute moral agency to AI and human decisions on service satisfaction. A quantitative, experimental design study revealed that in morally questionable service contexts, consumers' propensity to attribute morality is not influenced by the type of decision maker, AI or human. The ability of AI and humans to behave morally is rated similarly. The propensity to

attribute agency is significantly higher in the case of a human decision maker. The propensity to attribute moral agency is significantly positively affected by the favourability of the decision. A favourable decision increases moral agency, an unfavourable decision decreases moral agency, but this relationship is independent of whether the decision was made by a human or an AI. Perceived AI agency remained similar for favourable and unfavourable decisions, and lower than for humans. The tendency to attribute moral agency has a significant positive impact on the perceived fairness of the decision-maker. However, the type of decision-maker does not moderate this relationship. Perceived fairness mediates the relationship between the propensity to attribute morality and service satisfaction. However, perceived fairness does not mediate the relationship between the propensity to attribute agency and satisfaction. Whilst not hypothesised, a direct, positive relationship has been observed between the tendency to attribute morality and consumer satisfaction; in the case of agency, this relationship is negatively skewed.

Turinys

Lentelių sąrašas	9
Paveikslų sąrašas	11
Įvadas	12
1. Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis tyrimo aktualumas ir problematika	15
2. Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis teorinis pagrindimas	24
2.1. Moralinis veiksnumas ir jo reikšmė paslaugų kontekste	24
2.2. Polinkį priskirti moralinį veiksnumą DI ir žmogaus sprendimams paaiškinančios teorijos	30
2.3. Sprendimo palankumo ir polinkio priskirti moralinį veiksnumą ryšys	35
2.4. Suvokiamo sąžiningumo konceptualizacija	37
2.5. Suvokiamo sąžiningumo ir pasitenkinimo paslaugomis ryšys	43
2.6. Konceptualus vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis modelis.....	45
3. Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis tyrimo metodologija.....	50
3.1. Empirinio tyrimo objektas, tikslas ir uždaviniai.....	50
3.2. Empirinio tyrimo tipas, metodai ir tyrimo konstruktas	51
3.3. Empirinio tyrimo eigos ir duomenų analizės procedūros	55
4. Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis empirinio tyrimo rezultatai ir diskusija.....	60
4.1. Bendrosios tyrimo dalyvių sociodemografinės charakteristikos	60
4.2. Tyrimo instrumento metodologinės kokybės įvertinimas	63
4.3. Polinkio priskirti moralinį veiksnumą, suvokiamo sąžiningumo ir pasitenkinimo kintamųjų charakteristikos	64
4.4. Vartotojų polinkio priskirti moralinį veiksnumą, sprendimo palankumo, suvokiamo sąžiningumo ir pasitenkinimo sąsajų tyrimo rezultatų analizė	69
4.5. Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis empirinio tyrimo rezultatų diskusija	85
Išvados	90
Literatūros sąrašas	93
Informacijos šaltinių sąrašas	104
Priedai.....	105
1 priedas. SMV vertinimo veiksniai ir juos reprezentuojantys elementai	105
2 priedas. Suvokiamo proto veiksniai	106
3 priedas. Moralinių pagrindų klausimynas	107
4 priedas. Suvokiamo sąžiningumo ir vartotojų pasitenkinimo paslaugomis skalė	109
5 priedas. Tyrimo anketa	110
6 priedas. Tyrimo scenarijai	117
7 priedas. G*Power skaičiuotės rezultatai	119
8 priedas. Sociodemografinės tyrimo dalyvių charakteristikos.....	120
9 priedas. Faktorinė analizė	126
10 priedas. Aprašomoji statistika	138
11 priedas. Dėmesio ir manipuliacijos patikra	139

12 priedas. Raiškos pagal sociodemografines charakteristikas patikra.....	143
13 priedas. Homogeniškumo patikra.....	148
14 priedas. Kolmogorovo-Smirnov testas	154
15 priedas. Koreliacinė analizė.....	155
16 priedas. Tarpgrupinė dispersinė analizė (ANOVA)	156
17 priedas. Tiesinės paprastosios ir tiesinės daugialypės regresijos analizės	159

Lentelių sąrašas

1 lentelė. Suvokiamo DI moralinio veiksnio tyrimų apžvalga.....	17
2 lentelė. Vartotojų polinkio priskirti moralinį veiksnumą DI sprendimams tyrimų apžvalga	20
3 lentelė. Moraliniai pagrindai	33
4 lentelė. Suvokiamo DI sąžiningumo dimensijos ir taisyklės	39
5 lentelė. 2x2 tarpgrupinio dizaino schema.....	53
6 lentelė. Tyrimo konstrukto matavimo skalės lietuvių kalba	55
7 lentelė. Nestatistinio lyginamojo DI ir žmogaus sprendimų rinkodaroje vertinimo susijusioje tematykoje imties dydžio nustatymas.....	56
8 lentelė. Empirinio tyrimo imties dydžio nustatymas.....	57
9 lentelė. Dėmesio ir manipuliacijos patikros duomenys (N=409).....	60
10 lentelė. Demografinės tyrimo dalyvių charakteristikos (N=313).....	60
11 lentelė. Tyrimo dalyvių pasiskirstymas pagal šalį (N=313)	61
12 lentelė. Tyrimo dalyvių pasiskirstymas pagal amžių (N=313)	61
13 lentelė. Tyrimo dalyvių pasiskirstymas pagal amžiaus grupes (N=313)	61
14 lentelė. Respondentų pasiskirstymas pagal patirtį, naudojantis DI paremtais sprendimais paslaugose (N=313).....	62
15 lentelė. Respondentų pasiskirstymas pagal plaukų priežiūros paslaugų vartojimo dažnį (N=313)	62
16 lentelė. Respondentų pasiskirstymas pagal lytį, naudojantis plaukų priežiūros paslaugomis (N=313)	62
17 lentelė. Tiriamų konstrukto skalių tinkamumo vertinimas (N=313).....	63
18 lentelė. Tiriamų konstrukto skalių faktorinės analizės rezultatai (N=313).....	63
19 lentelė. Tiriamų konstrukto skalių patikimumo vertinimas (N=313)	64
20 lentelė. Tyrimo kintamųjų charakteristikos (N=313).....	65
21 lentelė. Sprendimo priėmimo ir sprendimo rezultato manipuliacijos patikros rezultatai (N=313).....	65
22 lentelė. Sprendimo palankumo ir sprendimo priimtumo vidurkių palyginimas (N=313)	66
23 lentelė. Scenarijaus realistiškumo analizė (N=313)	66
24 lentelė. Tyrimo kintamųjų raiškos priklausomybė nuo lyties (N=313)	67
25 lentelė. Kolmogorovo-Smirnovo testo rezultatai (N=313)	69
26 lentelė. Vartotojų polinkio priskirti moralinį veiksnumą, sprendimo palankumo, suvokiamo sąžiningumo ir pasitenkinimo kintamųjų koreliacinė analizė (N=313)	69
27 lentelė. Tyrimo dalyvių pasiskirstymas pagal scenarijus (N=313)	71
28 lentelė. Lygių dispersijų prielaidos Leveno testas: sprendimo priėmimo ir sprendimo palankumo poveikis vartotojų polinkiui priskirti moralinį veiksnumą (N=313)	71
29 lentelė. Sprendimo priėmimo ir sprendimo palankumo poveikis priklausomiems kintamiesiems (N=313)	71
30 lentelė. Tarpgrupinė nepriklausomų imčių dispersinė analizė: sprendimo priėmimo ir sprendimo palankumo poveikio polinkiui priskirti moralumą testas (N=313)	72
31 lentelė. Polinkio priskirti moralumą vidurkių palyginimai skirtingų scenarijų grupėse (N=313)	72
32 lentelė. Tarpgrupinė nepriklausomų imčių dispersinė analizė: sprendimo priėmimo ir sprendimo palankumo poveikio polinkiui priskirti veiksnumą testas (N=313)	74
33 lentelė. Polinkio priskirti veiksnumą vidurkių palyginimai skirtingų scenarijų grupėse (N=313)	75

34 lentelė. Daugialypės tiesinės regresijos, tiriant ryšį tarp polinkio priskirti moralinį veiksnumą ir suvokiamo sąžiningumo, priklausomo kintamojo rezultatai (N=313)	76
35 lentelė. Daugialypės tiesinės regresijos, tiriant ryšį tarp polinkio priskirti moralinį veiksnumą ir suvokiamo sąžiningumo, nepriklausomų kintamųjų rezultatai	77
36 lentelė. Regresijos modelių rezultatai, tiriant polinkio priskirti moralumą ir sprendimo priėmėjo sąveikos poveikį suvokiamam sąžiningumui (N=313).....	77
37 lentelė. Regresijos modelių rezultatai, tiriant polinkio priskirti veiksnumą ir sprendimo priėmėjo sąveikos poveikį suvokiamam sąžiningumui (N=313).....	78
38 lentelė. Tiesinės regresijos, tiriant ryšį tarp suvokiamo sąžiningumo ir pasitenkinimo, kaip priklausomo kintamojo rezultatai (N=313)	79
39 lentelė. Daugialypės tiesinės regresijos, tiriant ryšį tarp suvokiamo sąžiningumo ir pasitenkinimo, nepriklausomo kintamojo rezultatai	80
40 lentelė. Regresijos modelių rezultatai, tiriant sprendimo priėmėjo ir suvokiamo sąžiningumo poveikį pasitenkinimui (N=313).....	80
41 lentelė. Regresijos modelių rezultatai, tiriant ryšį tarp suvokiamo sąžiningumo, moralumo ir pasitenkinimo (N=313).....	81
42 lentelė. Tiesioginė, netiesioginė (per suvokiamą sąžiningumą) ir suminė moralumo įtaka pasitenkinimui (N=313).....	82
43 lentelė. Regresijos modelių rezultatai, tiriant ryšį tarp suvokiamo sąžiningumo, veiksnumo ir pasitenkinimo (N=313).....	83
44 lentelė. Tiesioginė, netiesioginė (per suvokiamą sąžiningumą) ir suminė veiksnumo įtaka pasitenkinimui (N=313).....	83
45 lentelė. Empirinio tyrimo hipotezių tikrinimo rezultatai.....	84

Paveikslų sąrašas

1 pav. Konceptualus vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikis pasitenkinimui paslaugomis modelis	49
2 pav. Eksperimento dizainas ir anketos struktūra.....	52
3 pav. Polinkio priskirti moralumą vidurkių palyginimas tarp nepriklausomų grupių.....	73
4 pav. Polinkio priskirti veiksnumą vidurkių palyginimas tarp nepriklausomų grupių.....	75
5 pav. Polinkio priskirti moralumą poveikis suvokiamam sąžiningumui, priklausomai nuo sprendimo priėmėjo tipo.....	78
6 pav. Polinkio priskirti veiksnumą poveikis suvokiamam sąžiningumui, priklausomai nuo sprendimo priėmėjo tipo.....	79
7 pav. Suvokiamo sąžiningumo poveikis pasitenkinimui paslauga, priklausomai nuo sprendimo priėmėjo tipo.....	81
8 pav. Moralumo įtaka pasitenkinimui paslauga per suvokiamą sąžiningumą.....	82
9 pav. Veiksnumo įtaka pasitenkinimui paslauga per suvokiamą sąžiningumą	84

Įvadas

Pasaulinės rinkos tendencijos rodo eksponentiškai augantį dirbtinio intelekto (DI) technologijų diegimą verslo sektoriuje. Naujausi duomenys atskleidžia, kad 2024 m., net 72% įmonių integravo DI į bent vieną verslo funkciją, o staigus šio rodiklio prieaugis rodo tendencingą reiškinio augimą (Statista, 2024). Teigiama, kad 2025 m., 89% įmonių didins arba išlaikys skiriamas investicijas į DI įvairiose industrijose (IBM, 2024). Be to, prognozuojama, jog ateinančiais metais kas trečia įmonė investuos daugiau nei 25 milijonus į DI, o daugiau nei 80% investicijų bus skirta verslo funkcijų pertvarkymui bei naujų paslaugų diegimui, siekiant padidinti veiklos efektyvumą ir konkurencingumą rinkoje (BCG, 2024). Įmonės, vertindamos DI strategijos efektyvumą, nurodo, jog vienas dažniausiai sekamų rodiklių yra vartotojų patirties ir pasitenkinimo gerinimas (IBM, 2024). Mokslininkai prognozuoja, jog DI sistemų diegimas paslaugų sektoriuje tęsis, siekiant pakeisti žmogiškuosius darbuotojus arba papildyti jų atliekamas užduotis daugelyje pramonės šakų (Pitardi et al., 2022). Šie plataus masto DI plėtojimo procesai, grindžiami teikiamų paslaugų, kurioms paprastai reikia žmogaus intelekto, automatizavimu, siekiant optimizuoti įmonių veiklą bei pagerinti sąveiką su vartotojais.

Didėjant investicijoms į DI technologijų paremtų sprendimų plėtrą, viena iš pagrindinių kliūčių, keliančių verslų susirūpinimą, yra su tuo susijusios etinės problemos (IBM, 2024). Vadybiniu požiūriu skiriamas didelis dėmesys suprasti ir numatyti galimą DI poveikį vartotojams ir visuomenei plačiąja prasme (Brendel et al., 2021). Anot Brendel ir kt. (2021), susirūpinimas etiniais iššūkiais yra itin svarbus tose verslo srityse, kuriose DI glaudžiai sąveikauja su žmonėmis, pavyzdžiui klientų aptarnavime. Plėtojant DI taikymą paslaugų sektoriuje, kinta paslaugų teikimo pobūdis, kuomet vartotojų ir paslaugą teikiančių žmonių sąveiką keičia DI sistemos. Didėjant DI pritaikomumui skirtinguose paslaugos teikimo etapuose, įskaitant savarankišką sprendimų priėmimą, kyla moralinių pasekmių rizika, kuri gali veikti vartotojų atsaką (Laakasuo et al., 2021). Nepaisant to, etiniu ir moraliniu požiūriu svarbiose paslaugų teikimo aplinkybėse, vartotojų ir DI sąveikos poveikis vartotojų elgsenai yra vis dar menkai nagrinėtas (Pitardi et al., 2022). Be to, trūksta žinių, kaip moraliniu požiūriu abejotinas sprendimas veikia vartotojų atsaką, kai jį priima paslaugą teikiantis žmogus ir tarp to egzistuojančius skirtumus, kai tai atlieka DI sistemos.

Tyrimo problema. Mokslininkų teigimu, žmonių sąveikos su DI priemonėmis supratimas yra svarbus, tačiau vis dar nepakankamai ištirtas reiškinys (Pitardi et al., 2022). Ankstesni tyrimai parodė, kad žmonės, sąveikaudami su DI sveikatos paslaugose, yra linkę atsisakyti DI priimamų sprendimų, kurie gali būti susiję su moraliniais klausimais (Longoni et al., 2019). Tuo tarpu Maninger'io ir Shank'o (2022) tyrimas parodė, kad moraliniu požiūriu svarbiose aplinkybėse, žmonės labiau pasitiki DI sprendimais nei žmogaus, o tam įtakos turi numanomas priimto DI sprendimo teisingumas. Teigiama, kad vartotojų reakcija į DI priklauso nuo moralinės situacijos kompleksiskumo ir kontekstinių aplinkybių. Haupt'as ir kt. (2024) pabrėžia, jog vartotojai, moraliniu požiūriu nevienareikšmiškai vertina DI taikymą versle ir rinkodaroje. Mokslininkai rekomenduoja plėsti rinkodaros tyrimų lauką, nagrinėjant, kaip vartotojų moralinis vertinimas veikia DI sistemų priėmimą arba atmetimą, nes tai turi tiesioginį poveikį tolimesniems įmonių vertinimams (Haupt et al., 2024). Garvey'is ir kt. (2023) pabrėžia, kad nepaisant dirbtinio intelekto galimybių paslaugų sektoriuje, trūksta žinių apie tai, kaip vartotojai vertina lygiaverčius DI ir žmogaus moralinius sprendimus. Mokslininkų atliktas tyrimas parodė, kad lyginami DI algoritmų ir žmonių pateiktus rinkodarinius pasiūlymus, moraliniu požiūriu vartotojai yra linkę vertinti DI sprendimus, kaip teisingesnius (Garvey et al., 2023).

Neetiški DI sprendimai vartotojų atžvilgiu, gali būti suvokiami, kaip moraliai neteisingi. Vartotojai, dalyvaudami neetiškose situacijose, yra linkę moraliai svarstyti ir vertinti su jais sąveikaujančius subjektus, įskaitant DI (Zhang et al., 2023). Šiam vertinimui įtakos turi suvokiamas subjekto veiksnumas (angl. *agency*) ir moralumas. Priskiriamo moralinio veiksnumo (angl. *moral agency*) reiškinys nagrinėtas psichologijos, sociologijos, dirbtinio intelekto ir etikos mokslų srityse. Tuo tarpu rinkodaros lauke, moralinio veiksnumo reiškinys, ypač DI ir vartotojų sąveikos kontekste, dar yra menkai nagrinėtas. Ankstesni tyrimai analizavo moralinio veiksnumo dedamąsias arba susijusius veiksnus, pavyzdžiui, žmonių polinkį vertinti priimamų sprendimų ar elgsenos moralumą, priskirti moralinę atsakomybę, moralinę kaltę, vertinti subjekto veiksnumą. Trūksta tyrimų, kurie nagrinėtų vartotojų polinkį priskirti moralinį veiksnumą DI ir žmogaus sprendimams paslaugų kontekste. Nėra tyrimų, kurie analizuotų, kaip priskiriamą moralinį veiksnumą veikia priimto sprendimo palankumas. Be to, nežinoma, kaip polinkis priskirti moralinį veiksnumą veikia suvokiamą sprendimo priėmėjo, DI ir žmogaus, sąžiningumą. Laakasuo ir kt. (2021) akcentuoja, kad nors suvokiamas priimtų sprendimų moralumas gali veikti vartotojų atsaką į paslaugas, galimas poveikis vartotojų elgsenai ir tai lemiančios priežastys yra mažai nagrinėtos. Vertinant DI arba kitų dirbtinių agentų veikimo įtaką vartotojams, neretai nagrinėjamas pasitikėjimo veiksnys (Söderlund, 2023a). Tačiau, trūksta žinių apie tai, kaip priskiriamas moralinis veiksnumas veikia vartotojų pasitenkinimą. Teigiama, jog vartotojų pasitenkinimas, lyginant su pasitikėjimu, turi didesnės įtakos teigiamiems vartotojų ketinimams ir su tuo susijusiems įmonių finansiniams rezultatams (Söderlund, 2023a). Esamas ištyrimo lygmuo nepateikia aktualių atsakymų apie tai, kaip priskiriamą moralinį veiksnumą veikia vartotojus aptarnaujantis subjektas (DI ar žmogus), subjekto priimto sprendimo palankumas ir kokią vaidmenį atlieka suvokiamas sąžiningumas, aiškinant ryšį tarp moralinio veiksnumo ir vartotojų pasitenkinimo paslaugomis. Todėl, atsižvelgiant į mokslinės literatūros spragas, formuluojami **problemniniai tyrimo klausimai**:

Kuriam paslaugų teikimo subjektui (DI ar žmogui) ir kokiomis aplinkybėmis vartotojai labiau linkę priskirti moralinį veiksnumą?

Kaip vartotojų polinkis priskirti moralinį veiksnumą DI ir žmogaus sprendimams veikia pasitenkinimą paslaugomis?

Tyrimo objektas – vartotojų priskiriamo moralinio veiksnumo DI ir žmogaus sprendimams poveikis pasitenkinimui paslaugomis.

Tyrimo tikslas – teoriškai ir empiriškai pagrįsti vartotojų polinkį priskirti moralinį veiksnumą DI ir žmogaus sprendimams ir jo poveikį pasitenkinimui paslaugomis.

Tyrimo uždaviniai:

1. Atskleisti vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis aktualumą ir problematiką;
2. Teoriškai argumentuoti konceptualų modelį, paaiškinantį vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikį pasitenkinimui paslaugomis;
3. Sudaryti vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis, tyrimo metodologiją;

4. Empiriškai patikrinti vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikį pasitenkinimui paslaugomis;
5. Pateikti vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis, praktines rekomendacijas ir tolimesnių tyrimų kryptis.

Tyrimo metodai. Sisteminė, lyginamoji mokslinės literatūros analizė, kiekybinis, tarpgrupinio eksperimento dizaino tyrimas, duomenų rinkimo metodas – anketinė apklausa. Statistinės duomenų analizės metodai: aprašomoji, faktorinė, neparametriniai testai, koreliacija, dispersija, paprastoji tiesinė ir daugialypė tiesinė regresija. Duomenys apdoroti naudojant *IBM SPSS* programinę įrangą.

1. Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis tyrimo aktualumas ir problematika

Dirbtinis intelektas paslaugose. Dirbtinis intelektas (DI) apibrėžiamas, kaip programų, algoritmų ir sistemų derinys, kuris „[...] imituoja žmogaus intelektą naudodamas algoritmus, pavyzdžiui, kontroliuojamą, nekontroliuojamą arba sustiprintą mašininių mokymąsi (ML), natūralios kalbos apdorojimą (NLP), gilųjų mokymąsi, robotų procesų automatizavimą ir taisyklėmis grindžiamas ekspertines sistemas.“ (Aker et al., 2023, p. 2). Anot mokslininkų Bogozzi ir kt. (2022), paslaugose DI gali būti pasitelkiamas įprastoms, mechaninio veikimo reikalaujančioms užduotims atlikti, analitinio pobūdžio užduotims arba užduotims, kurios reikalauja supratimo bei empatijos. Dėl savo techninių savybių, DI gali rinkti vaizdinius ir/ar tekstinius duomenis, juos analizuoti, mokytis iš jų ar teikti susijusias rekomendacijas bei kito pobūdžio grįžtamąjį ryšį vartotojams (Huang & Rust, 2021).

DI diegimas įvairiose paslaugų srityse sparčiai auga, siekiant padidinti klientų aptarnavimo efektyvumą ir pagerinti vartotojų patirtį. Mokslininkai teigia, kad tokios technologijos, kaip dirbtinis intelektas, iš esmės keičia paslaugų teikimo pobūdį, kuomet vartotojų ir paslaugą teikiančių žmonių sąveiką vis dažniau keičia DI sistemos, veikiant vartotojų patirtį ir santykį su paslaugomis (Xu et al., 2020). Nors sparti DI sprendimų plėtra įvairiose pramonės šakose lemia kintančias paslaugų teikimo sąlygas, mokslininkai pastebi, jog vis dar trūksta tyrimų, nagrinėjančių DI ir vartotojų sąveikos poveikį vartotojų elgsenai (Pitardi et al., 2022).

Siekiant nustatyti galimą DI sistemų ir vartotojų sąveikos paslaugose poveikį vartotojų atsakui į paslaugas, svarbu suprasti kaip vartotojai vertina DI moralinius sprendimus. Autoriai pabrėžia, jog būtent tai, kaip suvokiamos DI priemonės, priklauso jų elgsenos ir moralinių sprendimų vertinimas (Laakasuo et al., 2021). Tai, kaip žmonės suvokia DI sprendimų veiksnumą ir turimą emocinį lygį, reikšmingai lemia teigiamą arba neigiamą vartotojų reakciją (Pitardi et al., 2022). Nors vartotojų suvokiamas DI moralinis veiksnumas gali daryti įtaką paslaugos suvokimui ir vertinimui, vis dar mažai žinoma, kas daro įtaką jų, kaip moraliai veikiančių subjektų suvokimui (Laakasuo et al., 2021).

Siekiant identifikuoti svarbiausias tiriamojo lauko sritis ir trūkumus, būtina atlikti išsamią mokslinių tyrimų, nagrinėjančių vartotojų polinkio priskirti moralinį veiksnumą DI sprendimams, analizę. Nustačius su tyrimų lauku susijusias mokslinių tyrimų spragas ir galimas tolimesnes tyrimų kryptis, tyrimų aktualumo ir problematikos dalyje apžvelgiamos dvi mokslinių tyrimų sritys: suvokiamas DI moralinis veiksnumas ir vartotojų polinkis priskirti moralinį veiksnumą DI sprendimams.

Suvokiamo DI moralinio veiksnumo ištirtumo lygmuo. Mokslinės literatūros analizė atskleidė, jog šis reiškinys analizuojamas skirtingose mokslo srityse, kuriose nagrinėjami veiksniai ir aplinkybės, lemiančios žmonių suvokiamą DI moralinį veiksnumą. Ši mokslinė tema nagrinėjama psichologijos, sociologijos, dirbtinio intelekto ir etikos mokslų kryptyse, siekiant suprasti žmonių požiūrį į DI skirtinguose moraliniuose kontekstuose.

Mokslininkai Ayad'as ir Plaks'as (2025) savo tyrimu siekė ištirti žmonių priskiriamą moralinę atsakomybę DI bei nustatyti, kaip suvokiami socialiniai, ir psichologiniai faktoriai lemia šį polinkį moralinio nusizengimo atvejais. Analizuojama, kaip suvokiama DI autonomija, moralinės vertybės, emocinis savęs suvokimas ir socialiniai ryšiai daro įtaką priimtų sprendimų moralumui. Tyrėjai nagrinėjo, kaip vartotojai vertina DI moralinius sprendimus, kai jie buvo suvokiami, kaip tyčiniai ar atsitiktiniai. Rezultatai atskleidė, jog žmonių polinkiui priskirti moralinę atsakomybę DI, įtakos turi suvokiamas DI (ne)sąmoningumas. Kai suvokiama, jog DI atlikti veiksmai yra labiau tyčiniai nei

atsitiktiniai, žmonės yra linkę juos suvokti, kaip moraliai neteisingus. Ši išvada rodo, jog tokiu atveju, žmonės yra linkę svarstyti DI, kaip galintį savarankiškai priimti moralinius sprendimus. Žmonės suvokia DI, kaip nepriklausomą subjektą, kuris gali priimti moralinę atsakomybę. Tuo tarpu netyčinių moralinių nusižengimų atveju, ypač, kai DI pasižymėjo socialiniais ryšiais, jam buvo priskirtas silpnesnis moralinis veiksnumas ir didesnė priklausomybė išoriniams veiksniams.

Autoriai teigia, jog moralinės situacijos, kuriose žmogus sąveikauja su DI, neretai apima įvairius žmogaus vaidmenis ir bendrus DI ir žmogaus sprendimus (Dong & Bocian, 2024). Todėl, Zhang ir kt. (2023) moksliniu tyrimu siekta nustatyti, kaip skirtingų tipų moralinės dilemos veikia žmonių suvokiamą DI ir žmogaus moralinį veiksnumą. Trys eksperimentai atskleidė, jog tai priklauso nuo situacijos konteksto, kuomet skirtingų tipų moralinės dilemos veikiamos skirtingų psichologinių ir neurologinių mechanizmų. Jei moralinė situacija kelia netiesioginę žalą, žmonės DI sistemas gali vertinti griežčiau nei žmones, vertinti DI veiksmus, kaip mažiau priimtinius. Itin sudėtingose moralinėse dilemose, žmonės tikisi, kad DI atliks moraliai teisingus sprendimus ir yra mažiau atlaidesni, kai šie lūkesčiai nepateisinami. Tačiau, tiesioginės moralinės žalos atveju, dėmesys krypta į paties atlikto veiksmo moralumą, tuomet tiek DI, tiek žmogaus moraliniai sprendimai vertinami lygiavertiškai.

Gamez'o ir kt. (2020) moksliniame darbe nagrinėjamas žmonių polinkis vertinti dirbtinius moralinius agentus (DMA) (angl. *artificial moral agents*, AMA), įskaitant DI, kaip moralinį subjektą ir jo veiksmams priskirti dorybę ar ydingumą, lyginant su žmonėmis. Eksperimentinio tyrimo rezultatai atskleidė, kad etiniu požiūriu, žmonės panašiai vertina dirbtinių moralinių agentų ir žmonių elgesį, tačiau suvokia, kad DMA, įskaitant DI, negali turėti moralinių ketinimų bei priimti moralinių sprendimų. Todėl, žmonės yra linkę manyti, kad DI negali būti moraliai veiksnus.

Maninger'io ir Shank'o (2022) moksliniame tyrime nagrinėjamas žmonių suvokiamas priimtų moralinių sprendimų (ne)teisingumas, lyginant DI ir žmonių moralinius pažeidimus, išskiriami moralinių nusižengimų vertinimo kriterijai bei tiriamas polinkis priskirti kaltę DI ir žmogui. Darbe tiriamos skirtingos moralinių nusižengimų ribinės situacijos, vertinant sukeltą žalą, priimtų veiksmų neteisingumą, išdavystę, nuvertinimą, pažeminimą bei priespaudą. Eksperimentai atskleidė, kad žmonės yra linkę vertinti DI priimtus sprendimus, kaip labiau teisingus, lyginant su identiškais žmonių priimtais moraliniais sprendimais. Be to, žmonės dažniau priskiria kaltę žmogui, dėl suvokiamo DI neveiksnumo veikti inicijuotai. Ši sąlyga kinta tiesioginės žalos atveju, kuomet kaltę priskiriama abiem subjektams vienodai.

Banks (2021) savo tyrimu siekė nustatyti, kaip suvokiamą DI grįsto roboto arba žmogaus moralinį veiksnumą veikia priimti moraliniai sprendimai ir elgsena, kuri pažeidžia pagrindinius moralinius principus, kaip rūpestingumas, teisingumas ir t.t.. Nagrinėjama, kaip priskiriamas moralumas, atsakomybė ir pasitikėjimas kinta priklausomai nuo sąveikos pobūdžio (gyvo ar stebint) bei subjekto tipo – žmogaus arba DI roboto. Nustatyta, kad žmonės yra linkę panašiai vertinti DI paremto roboto ir žmogaus moralinę elgseną, tačiau robotui priėmus moraliai neteisingą sprendimą, jis vertinamas griežčiau. Be to, suvokiamam moraliniam veiksnumui įtakos turi tiesioginė sąveika. Būtent ši sąlyga lėmė tai, jog DI robotui buvo priskirta moralinė atsakomybė.

Žemiau pateiktoje lentelėje, pristatomi minėtų tyrimų autoriai, tyrimų objektai ir tikslai, pateikiami gauti tyrimo rezultatai, rekomenduojamos tolimesnių tyrimų kryptis (žr. **1 lentelė**).

1 lentelė. Suvokiamo DI moralinio veiksnio tyrimų apžvalga

Autorius (-iai), metai	Tyrimo objektas	Tyrimo tikslas	Tyrimo rezultatai	Tolimesnių tyrimų kryptys
Ayad ir Plaks (2025)	Žmonių polinkis priskirti moralinę atsakomybę DI, šį poveikį lemiantys psichologiniai ir socialiniai veiksniai.	Atsižvelgiant į tyčinius ir netyčinius moralinius pažeidimus, įvertinti žmonių polinkį priskirti moralinę atsakomybę DI ir kokią įtaką tam daro psichologiniai bei socialiniai veiksniai.	Tyrimas parodė, kad žmonės linkę priskirti moralinį veiksnumą DI sprendimams , kurie buvo atlikti tyčia ir sukėlė neigiamas pasekmes . Šiam vertinimui įtakos turėjo psichologinės priežastys, siejamos su polinkiu priskirti veiksnumą – suvokiamas DI autonomiškumas, savarankiškumas, gebėjimas veikti turint ketinimų.	Toliau nagrinėti, kaip atskirų psichologinių faktorių skirtumai daro įtaką žmonių moraliniams vertinimams. Ištirti, kaip skiriasi suvokiamas DI moralumas dviprasmiškose moralinėse aplinkybėse . Tyrimuose įtraukti žmonės, siekiant įvertinti skirtumus tarp DI ir žmogaus .
Zhang ir kt. (2023)	DI ir žmonių sprendimų moralinis vertinimas skirtingose moralinėse dilemose.	Palyginti, kaip žmonės vertina DI ir žmonių priimtus moralinius sprendimus moralinėse dilemose, ir nustatyti, kokias moralines normas žmonės yra linkę priskirti dirbtinio intelekto sistemoms.	Rezultatai atskleidė, jog žmonės skirtingai vertina DI priimtus moralinius sprendimus, skirtingose moralinėse dilemose, o šiam poveikiui įtakos turi kontekstinė aplinka, kurios suvokimą lemia skirtingi situacijos vertinimo mechanizmai. Tyrimas parodė, kad moraliniu požiūriu mažiau sudėtingose aplinkybėse, žmonės yra linkę nuvertinti DI moralinius sprendimus ir juos vertinti, kaip mažiau teisingus , lyginant su identiškais žmonių sprendimais. Itin sudėtingose moralinėse situacijose, abiem subjektams taikomos vienodos moralės normos .	Toliau tirti priežastis, kurios lemia žmonių polinkį skirtingai vertinti DI moralinius gebėjimus, priklausomai nuo moralinės dilemos poveikio. Nustatyti, kas lemia neigiamą atsaką į DI moralinius sprendimus, kodėl, kai kuriais atvejais, taikomos tos pačios moralės normos tiek žmogui, tiek DI .
Gamez ir kt. (2020)	Moralinio veiksnio priskyrimas DI ir jo moralinių sprendimų vertinimas (ydingumo ir dorybės atžvilgiu), lyginant su žmonėmis.	Ištirti žmonių polinkį vertinti DI, kaip moralų subjektą, kurio sprendimai gali būti vertinami ydingai arba dorai, lyginant su identiškais žmogaus moraliniais sprendimais.	Nustatyta, kad žmonės nėra linkę vertinti DI, kaip moralaus subjekto dėl numanomo negebėjimo turėti emocijų, moralinių ketinimų ir priimti moralinių sprendimų. Todėl, palyginti su žmogumi, DI moralinis veiksnumas yra silpnesnis.	Rekomenduojama plėsti dirbtinių moralinių agentų ir žmogaus moralinių sprendimų vertinimo aplinkybes ir kontekstus, siekiant geriau suprasti ar tie patys atributai, kurie priskiriami žmonėms mažina suvokiamą žmonių kaltę, yra vienodi DI atveju.
Maninger ir Shank (2022)	Moralinių nusižengimų vertinimas, priskiriant kaltę DI ir žmogui.	Palyginti DI ir žmonių moralinius sprendimus moralinių nusižengimų situacijose ir nustatyti, kam žmonės yra linkę priskirti kaltę.	Tyrimas parodė, kad žmonių moraliniai nusižengimai, apimantys žalos, išdavystės ir nuvertinimo atvejus, vertinami prasčiau nei DI atžvilgiu. Žmonės yra linkę labiau pasitikėti DI moraliniais sprendimais , dėl numanomo negebėjimo pažeisti iš anksto numatytą, „užprogramuotą“ teisingąjį pagrindą. Nustatyta, kad žmonės	Siūloma įvertinti tarpinius veiksnus, kurie gali turėti įtakos suvokiamiems DI ir žmogaus moraliniams sprendimams, pavyzdžiui emocinė būseną. Įvertinti, kaip kinta žmonių reakcija, priimant moraliai gerus sprendimus. Siūloma nagrinėti, kokią įtaką turi fizinis DI įkūnijimas.

			laikomi linkę dažniau nusizengti ir klysti moraliniuose sprendimuose nei DI sistemos, išskyrus žalos atvejus, kuomet kaltės priskyrimas yra lygiavertis abiejų subjektų atžvilgiu.	Įvertinus elgsenos ir fizinius pokyčius, nustatyti ar tai turi įtakos žmonių polinkiui priskirti moralinį veiksnumą.
Banks (2021)	DI paremto roboto ir žmogaus moralinio veiksnumo vertinimas ir poveikis atsakomybei, suvokiamam moralumui ir pasitikėjimui.	Palyginti, kaip roboto ir žmogaus priimti sprendimai bei elgsena moralinėse situacijose, veikia suvokiamą moralinį veiksnumą bei nustatyti, kaip tai lemia polinkį priskirti atsakomybę, moralumą ir pasitikėjimą.	Tyrimas parodė, jog žmonės DI ir žmogaus elgseną moraliniu požiūriu vertina panašiai . Tačiau, vertinant priimtų neteisingus moralinius sprendimus, griežčiau vertinami robotai . Nustatyta, kad tiesioginė sąveika turi įtakos suvokiamam roboto moraliniam veiksnumui , kai žmonės yra linkę priskirti didesnę moralinį veiksnumą nei stebint iš išorės, pavyzdžiui vaizdo įrašė.	Įvertinti, kaip kinta suvokiamas moralumas gero ir blogo elgesio atvejais . Nagrinėti scenarijus, kuriuose tyrimo dalyviai patys nukentėtų arba gautų naudos dėl DI ir žmogaus moralinių sprendimų . Atskirti fizinio artumo ir tiesioginės sąveikos aplinkybes, siekiant nustatyti ar suvokiamam moraliniam veiksnumui įtakos turi šių aplinkybių derinys, ar vienas iš jų.

Mokslinių tyrimų, nagrinėjančių žmonių polinkį priskirti moralinį veiksnumą DI analizė leido nustatyti, jog šia tema egzistuoja nevienareikšmės mokslinių tyrimų išvados. Autoriai teigia, kad žmonės gali būti linkę priskirti moralinį veiksnumą DI, o tai priklauso nuo kontekstinių aplinkybių, pavyzdžiui moralinių dilemų sudėtingumo (Zhang et al., 2023), sąveikos su DI pobūdžio (Banks, 2021) ir psichologinių veiksnių (Ayad & Plaks, 2025). Kita vertus, teigiama, kad žmonės gali nepriskirti moralinio veiksnumo, dėl suvokiamos DI autonomiškumo stokos, priimant moralinius sprendimus (Ayad & Plaks, 2025) bei nesavarankiškumo (Maninger & Shank, 2022). Taip pat, minėti veiksniai turi įtakos tam, kaip žmonės suvokia DI sprendimų moralumą, lyginant su žmogumi. Vieni autoriai teigia, kad, dėl numanomo DI neveiksnumo, DI sprendimai vertinami priimtinau nei žmogaus (Maninger & Shank, 2022), kiti teigia atvirkščiai – nepalankiau, arba panašiai (Zhang et al., 2023). Be to, tyrimuose pastebimos ir kraštutinės reikšmių vertės, kuomet teigiama, kad DI iš esmės nevertinamas, kaip moralus subjektas (Gamez et al., 2020).

Dviprasmiški mokslinių tyrimų rezultatai atskleidžia, jog egzistuoja įvairūs požiūriai, kuriuos dažniausiai lemia skirtingi kontekstiniai ir psichologiniai veiksniai. Nėra pakankamai žinių, kurios tiksliai apibrėžtų, kaip ir kodėl žmonės (ne)priskiria moralinio veiksnumo DI.

Vartotojų polinkio priskirti moralinį veiksnumą DI sprendimams ištirtumo lygmuo. Pastebima, jog rinkodaros mokslo literatūroje egzistuoja nedaug tyrimų, nagrinėjančių vartotojų polinkį priskirti moralinį veiksnumą DI sprendimams paslaugose. Todėl, šioje darbo dalyje pateikiami susijusioje tematikoje vykdytų tyrimų apžvalga rinkodaros mokslo srityje.

Atlikus mokslinės literatūros analizę, pastebėta, jog daugelis tyrėjų savo darbuose nagrinėjo atskiras moralinio veiksnumo dedamąsias, tokias, kaip vartotojų suvokiamą DI veiksnumą, vartotojų polinkį priskirti kaltę, atsakomybę ar savanaudiškumą DI atžvilgiu. Siekdami pagrįsti žmonių polinkį priskirti minėtus aspektus, autoriai išskiria susijusias subjektų grupes. Dažniausiai, šios grupės apima patį DI sprendimą, DI sprendimo kūrėjus, žmogiškuosius darbuotojus bei paslaugą teikiančią įmonę.

Sullivan ir Wamba (2022) moksliniame tyrime nagrinėjama žmonių priskiriama moralinė kaltė bei atsakomybė, kai paslaugos teikimo nesėkmė susijusi su DI sprendimo moraliniais pažeidimais.

Moksliniu darbu siekta nustatyti, kaip vartotojai suvokia, kas yra atsakingas už padarytą žalą, vertinant skirtingas puses – paslaugą teikiančią įmonę, DI sprendimo kūrėjus bei patį DI sprendimą. Trys eksperimentiniai tyrimai atskleidė, jog dažniausiai žmonės priskiria kaltę DI sprendimo kūrėjams, vėliau paslaugą teikiančioms įmonėms ir galiausiai pačiai DI technologija paremtai sistemai. Be to, tyrimas parodė, jog vartotojų priskiriamą atsakomybę lemia žalos pobūdžio suvokimas ir vertinimas. Jei sukelta žala vertinama labiau kaip tyčinė nei atsitiktinė, tuomet moralinė atsakomybė priskiriama DI sprendimui. Visa tai grindžiama vartotojų suvokiamais DI proto gebėjimais (angl. *mind perception*), t.y. gebėjimu racionaliai mąstyti, priimti sąmoningus sprendimus ir turėti ketinimų. Vis dėlto, didesnę įtaką priskiriamai kaltei ir atsakomybei DI atžvilgiu, turėjo suvokiama DI patirtis (gebėjimas jausti emocijas). Be to, priskiriamai atsakomybei tyčinės ar atsitiktinės žalos atveju įtakos turėjo ir žalos mastas. Jei žala buvo vertinama, kaip vidutinė ir nesukelianti itin skaudžių, negrįžtamų pasekmių, atsakomybė priskiriama labiau tyčinės žalos atveju nei atsitiktinės (Sullivan & Wamba, 2022).

Mokslininkų Garvey'o ir kt. (2023) atliktu tyrimu siekta praplėsti, anot jų, mažai tyrinėtą mokslinių tyrimų sritį ir išanalizuoti kaip vartotojai vertina DI sprendimus, palyginti su lygiaverčiais žmonių sprendimais. Tyrimu buvo siekiama įvertinti vartotojų polinkį priskirti savanaudiškų ir/ar nesavanaudiškų ketinimų tikimybę DI, teikiant rinkodarinius pasiūlymus ir, kaip tai veikia vartotojų pasitenkinimą bei sprendimus pirkti, lyginant su lygiaverčiais žmogaus pateiktais pasiūlymais. Eksperimentinio tyrimo rezultatai parodė, jog vartotojai mano, kad DI negali turėti moralinių ketinimų, priimti inicijuotų sprendimų ir nuspręsti ar suteiktas pasiūlymas bus savanaudiškas, ar geranoriškas kliento atžvilgiu. Be to, buvo nustatyta, kad lygiaverčių pasiūlymų vertinimas priklauso nuo priskiriamų ketinimų – savanaudiškumo ir geranoriškumo, DI sprendimui arba žmogui. Vertinant DI pasiūlymus, vartotojai mano, kad DI yra mažiau savanaudiškas, tuo tarpu žmogaus pasiūlymai vertinami, kaip labiau geranoriški.

Mokslininkų Arian ir kt. (2023) tyrimu siekta nustatyti vartotojų polinkį priskirti atsakomybę paslaugą teikiančiam DI robotui, sužmogintam DI robotui ir žmogui, negrįžtamo paslaugos gedimo bei paslaugos atkūrimo atvejais. Tyrimas parodė, jog įvykus paslaugos gedimui, vartotojai yra linkę priskirti didesnę veiksnumą žmogui, ypač, kai paslaugą teikia žmogiškų savybių neturintis robotas. Visa tai yra lemiamą vartotojų manymu, kad tokios užduotys, kaip paslaugų gedimų atkūrimas, yra sudėtingos emociniu ir socialiniu lygmeniu, todėl jos gali būti atliekamos tik žmonių. DI suvokiami, kaip negalintys savarankiškai sukelti gedimo ar išspręsti susijusių padarinių, nes jie nepasižymi veiksnumo savybėmis – jie negali veikti autonomiškai ir privalo laikytis nurodytų komandų (Arian et al., 2023). Todėl, tokiu atveju, vartotojai yra linkę priskirti atsakomybę žmogui.

Ryoo'as ir kt. (2024) savo moksliniame darbe nagrinėjo vartotojų polinkį priskirti kaltę paslaugą teikiančiam DI sprendimui, žmogui ar paslaugą atstovaujančiai įmonei, įvykus nelaimingam atsitikimui. Atlikti keturi tyrimai parodė, jog paslaugos teikimo metu, įvykus nelaimingam atsitikimui, dėl žmogaus kaltės, vartotojai atsakomybę priskirs žmogui. Tuo tarpu, jei žala bus sukelta paslaugų roboto, atsakomybė bus priskirta paslaugą teikiančiai įmonei. Tai grindžiama vartotojų nuomone, jog įmonė kontroliuoja paslaugą teikiančius robotus, tačiau atsakyti už paslaugą teikiančius žmones – negali. Buvo nustatyta, jog DI roboto intelektualiniai sugebėjimai lemia vartotojų priskiriama kaltę ir pasitenkinimą, t.y. kuo aukštesnis roboto intelektas, tuo didesnė tikimybė, jog vartotojai kaltins paslaugų robotą, o ne atstovaujančią įmonę ir, tikėtina, jog bus linkę grįžti ir pasinaudoti suteikta paslauga pakartotinai.

Pavone‘as ir kt. (2023) analizavo neigiamos vartotojų sąveikos su DI pokalbių robotu poveikį klientų neigiamoms emocijoms bei polinkiui priskirti atsakomybę dėl paslaugos nesėkmės, lyginant sąveiką su žmogumi. Tyrimas parodė, kad paslaugos nesėkmės atveju, vartotojai nėra linkę priskirti atsakomybės DI dėl numanomo negebėjimo kontroliuoti savo veiksmų, t.y. veikti inicijuotai ir pasižymėti savikontrole. Todėl, dėl neigiamo paslaugos teikimo rezultato, atsakomybė priskiriama įmonei. Pastebėta, jog DI suteikus žmogiškų savybių, sumažėja polinkis priskirti atsakomybę įmonei.

Žemiau pateiktoje lentelėje (žr. **2 lentelė**), pristatomi susijusioje tematikoje atliktų tyrimų apžvalga. Įvardijami tyrimų autoriai, apžvelgiami tyrimų objektai, hipotezės/tikslai, rezultatai bei pristatomos tyrėjų rekomenduojamos tolimesnių tyrimų kryptys.

2 lentelė. Vartotojų polinkio priskirti moralinį veiksnumą DI sprendimams tyrimų apžvalga

Autorius (-iai), metai	Tyrimo objektas	Tyrimo hipotezė (-ės) / tikslas	Tyrimo rezultatai	Tolimesnių tyrimų kryptys
Sullivan ir Wamba (2022)	Moralinės kaltės ir atsakomybės priskyrimui DI sistemai, DI sistemos kūrėjui arba paslaugą teikiančiai įmonei, DI sukeltos moralinės žalos atveju.	Atlikti trys tyrimai ir iškeltos keturios hipotezės. H1: Žmonės labiau kaltins DI, kai pažeidimas bus laikomas tyčiniu, nei kai jis bus laikomas atsitiktiniu. H2a: Suvoktas DI veiksnumas yra tarpininkas tarp suvokiamos tyčinės žalos (skirtos žmonėms) ir kaltės vertinimo DI atžvilgiu. H2b: Suvokiama DI patirtis (angl. <i>experience</i>) lemia ryšį tarp suvokiamos tyčinės žalos (nukreiptos į žmones) ir kaltės vertinimo DI atžvilgiu. H3: Žmonės labiau kaltins įmones, kai su DI susijęs pažeidimas bus suvokiamas kaip tyčinis, nei kaip atsitiktinis. H4: Žmonės labiau kaltins DI kūrėjus, kai su DI susijęs pažeidimas bus suvokiamas kaip tyčinis, kaip atsitiktinis.	Tyrimas atskleidė, jog dažniausiai žmonės kaltę priskiria DI sprendimo kūrėjams, vėliau paslaugą teikiančioms įmonėms ir galiausiai pačiai DI sistemai. Polinkiui priskirti kaltę ir atsakomybę DI, turėjo tyčinis žalos pobūdis . Šiuo atveju padidėjo suvokiamas DI veiksnumas (gebėjimas savarankiškai veikti, sąmoningai mąstyti). Vis dėlto, lyginant suvokiamą DI veiksnumą ir patirties (gebėjimo jausti emocijas) įtaką priskiriamai kaltei ir atsakomybei DI atžvilgiu, didesnę įtaką turėjo suvokiama DI patirtis.	Ištirti suvokiamo proto (angl. <i>mind perception</i>) skirtumų tarp skirtingų subjektų (pvz.: žmogaus ir DI) poveikį priskiriamai moralinei atsakomybei. Nagrinėti DI išvaizdos, lyties poveikį suvokiamai moralinei kaltei ir atsakomybei.
Garvey‘is ir kt. (2023)	Pasiūlymą teikiančio DI sprendimo ar žmogaus įtaka vartotojų reakcijai į neigiamai ir teigiamai vertinamus rinkodaros pasiūlymus, įvertinant subjektų moralinius ketinimus.	Atlikti 5 tyrimai ir iškeltos 3 hipotezės. H1: Pasiūlymo kaina, geresnė / blogesnė nei tikėtasi, mažina / didina vartotojo pasitenkinimą bei tikimybę įsigyti, kai šį pasiūlymą teikia DI sistema, lyginant su žmogumi. H2: H1 lemia vartotojų priskiriami moraliniai ketinimai – (ne)savanaudiškumas ir (ne)geranoriškumas DI ir žmogaus atžvilgiu. H3: DI sprendimui suteikus žmogiškų savybių, vartotojų reakcija į blogesnius nei tikėtasi pasiūlymus prastėja, o į geresnius nei tikėtasi gerėja.	Tyrimo rezultatai parodė, jog vartotojai dažniau priima prastesnius (nei tikėtasi) DI pasiūlymus, lyginant su žmogaus pateiktu identišku pasiūlymu. Šį poveikį lemia suvokiamas DI nesavanaudiškumas. Vartotojai mano, kad DI negali turėti inicijuotų ketinimų ir nuspręsti , kaip pateiktas pasiūlymas bus vertinamas iš kliento perspektyvos. Dėl šios priežasties DI sprendimai suvokiami, kaip labiau moralūs, nesavanaudiški.	Plėsti tyrimus, nagrinėjant, kaip vartotojai suvokia sprendimo sąžiningumą, atsižvelgiant į subjekto tipą – DI ar žmogų. Įvertinti, kaip suvokiamas sąžiningumas daro įtaką sprendimo priėmimui ir ketinimams . Įvertinti, kaip teigiamos ir neigiamos

				sąveikos aplinkybės lemia polinkį rinktis DI ar žmogaus sprendimus.
Arikan'as ir kt. (2023)	Vartotojų polinkis priskirti atsakomybę, dėl paslaugos nesėkmės, atkūrimo, paslaugą teikiančiam DI robotui, sužmogintam DI robotui ir žmogui.	Atlikti du eksperimentai ir iškeltos keturios hipotezės. H1: Kai nėra galimas paslaugos atkūrimas, paslaugos nesėkmė priskiriama žmogiškų savybių neturinčiam robotui, lyginant su žmogiškų savybių turinčiu robotu ir žmogumi. H2: Kai galimas paslaugos atkūrimas, atsakomybė už paslaugos atkūrimą bus priskiriama įmonei, kai paslaugą teiks žmogiškų savybių neturintis robotas, lyginant su žmogiškų savybių turinčiu robotu ir žmogumi. H3: Paslaugą teikiančio subjekto (roboto, sužmoginto roboto ar žmogaus) tipo poveikis vartotojų pasitenkinimui yra lemiamas vartotojų priskiriamu atlidumu. H3a: Kai paslaugos nesėkmė priskiriama įmonei, mažėja vartotojų atlidumas ir pasitenkinimas.	Tyrimas atskleidė, kad vartotojų polinkiui priskirti atsakomybę įtakos turi paslaugos teikimo subjektas. Dėl paslaugos nesėkmės, kai ją teikė žmogus, atsakomybė bus priskirta žmogui. Tuo tarpu paslaugos nesėkmės atveju, kai ją teikė DI, atsakomybė bus priskirta įmonei. Kai atsakomybė priskiriama įmonei, vartotojai yra linkę būti mažiau nuolaidesni, tikėdami, jog nesėkmė gali pasikartoti. Atsakomybė nėra priskiriama DI, dėl suvokiamo DI neveiksnumo.	Nustatyti, kaip kinta atsakomybės priskyrimas, kai paslaugos nesėkmė susijusi su paslaugų roboto sprendimais. Plėsti hipotezių taikymą skirtinguose paslaugų ir kultūriniuose kontekstuose.
Ryoo'as ir kt. (2024)	Vartotojų polinkis priskirti kaltę, priklausomai nuo paslaugos subjekto tipo – DI robotui, žmogui ar paslaugą teikiančiai įmonei.	Atlikti 4 tyrimai ir iškeltos 8 hipotezės, atsižvelgiant į jų aktualumą, žemiau pateikiamos 5, šio tyrimo kontekstui aktualios hipotezės. H1a: Vartotojai bus linkę priskirti mažesnę atsakomybę paslaugą teikiančiam robotui nei paslaugą teikiančiam žmogui. H1b: Mažesnę priskiriama atsakomybę bus lemiamas vartotojų suvokimu, jog paslaugą teikiantys robotai yra mažiau kontroliuojami, lyginant su paslaugą teikiančiais žmonėmis. H2a: Vartotojų priskiriamos kaltės mastas, paslaugą teikiančiai įmonei, bus lemiamas tuo, kad už padarytą žalą bus atsakingi paslaugą teikiantys robotai, o ne žmonės. H3b: Priskiriama kaltė bus lemiamas vartotojų suvokimu, jog į žmogų panašesnis robotas yra labiau kontroliuojamas nei žmogiškų savybių neturintis robotas. H4b: Roboto tipas, analitinis ar mechaninis ir suvokiama didesnė/mažesnė kontrolė lems poveikį.	Atliktas tyrimas parodė, jog paslaugos teikimo metu įvykus nelaimingam atsitikimui dėl žmogaus kaltės, vartotojai atsakomybę priskirs žmogui. Kita vertus, jei žala bus sukelta DI paslaugų roboto, atsakomybė bus priskirta paslaugą teikiančiai įmonei. Ištyrus, kaip kinta vartotojų suvokimas, kai robotui suteikiamos žmogiškosios savybės, nustatyta, jog priskiriama kaltė dėl klaidos dažniau tenka robotui, o ne paslaugą teikiančiai įmonei. Todėl, neigiamas poveikis susilpnintas, kai įmonės DI paslaugų robotams suteikia žmogiškųjų savybių. Priskiriamai atsakomybei bei kaltei įtakos turi roboto išvaizda ir intelektas. Kuo robotas panašesnis į žmogų, tuo padidėja vartotojų priskiriama kaltė ir atsakomybė DI robotui , o ne įmonei. Taip pat, aukštesnis roboto intelektas lemia šį poveikį.	Toliau nagrinėti, kaip kaltės priskyrimas DI robotui, žmogui ar įmonei, veikia vartotojų atsaką ilgalaikėje perspektyvoje, pavyzdžiui lojalumą. Ištirti, kokią įtaką roboto tipas turi vartotojų suvokiamai roboto kontrolei ir polinkiui priskirti kaltę.

Pavone, Meyer-Waarden ir Munzel (2023)	Su DI pokalbių robotu susijusios paslaugos teikimo nesėkmės poveikis vartotojų emocijoms, atsakomybės priskyrimui ir situacijos įveikimui.	Tyrėjai kelia 10 hipotezių, tačiau apžvalgoje pateikiamos tyrimui aktualiausios 4 hipotezės. H1: Klientai bus labiau linkę priskirti atsakomybę dėl paslaugos nesėkmės žmogui, o ne robotui. H2: Vartotojai bus labiau linkę priskirti atsakomybę, dėl paslaugos nesėkmės, įmonei, jei paslaugą teiks robotas, o ne žmogus. H9: Ketinimai (a) mažina atsakomybės priskyrimą įmonei, ir (b) šį poveikį lemia pokalbių roboto antropomorfinių vaizdinių užuominų lygis, t. y. ketinimai labiau mažina atsakomybės priskyrimą esant aukštam antropomorfinių vaizdinių užuominų lygiui nei žemam. H10: Antropomorfiniai vaizdiniai ženklai mažina atsakomybės priskyrimą įmonei.	Tyrimo rezultatai atskleidė, jog paslaugos nesėkmės atveju, žmonės nėra linkę priskirti atsakomybės DI pokalbių robotui, dėl suvokiamų ribotų proto gebėjimų – veikti turint išankstinių ketinimų ir kontroliuoti savo veiksmus. Tokiu atveju, pastebimas didesnis žmonių polinkis priskirti atsakomybę įmonei. Tuo tarpu, sąveikoje su žmogumi, atsakomybė priskiriama žmogui ir įmonei. Pokalbių robotui suteikus žmogiškų savybių, mažėja neigiamas vartotojų reakcija, požiūris į įmonę, padedant palankiau reaguoti į stresines paslaugos situacijas.	Ištirti, kitas susijusias vartotojų emocines reakcijas , taip pat nustatyti, kaip neigiamos paslaugos aplinkybės veikia vartotojų gerovę ir pasitenkinimą .
--	--	--	---	---

Susijusioje tematikoje vykdytų tyrimų apžvalga rodo, jog rinkodaros mokslo literatūroje, pakankamai siaurai nagrinėtas vartotojų polinkis priskirti moralinį veiksnumą DI, analizuojant tik atskirus šio reiškinio veiksnus. Autoriai nagrinėjo vartotojų polinkį priskirti atsakomybę (Arikan et al., 2023; Pavone et al., 2023), kaltę (Ryoo et al., 2024), arba atsakomybę ir kaltę kartu (Sullivan & Wamba, 2022). Taip pat, nagrinėtas vartotojų požiūris į DI arba žmogaus sprendimus, kurie lemiami suvokiamu DI veiksmumu (Garvey ir kt., 2023). Didžioji dalis tyrimų analizavo, kaip polinkį priskirti atsakomybę ar kaltę veikia suvokiamas DI protas, kuris yra vienas iš suvokiamo veiksnumo sąlygų. Tyrimai atskleidė, jog dažniausiai vartotojai nėra linkę priskirti atsakomybės ar kaltės DI, dėl suvokiamo DI neveiksnumo – negebėjimo racionaliai mąstyti, turėti ketinimų ar veikti autonomiškai. Tačiau egzistuoja ir priešingos sąlygos, kuomet vartotojų polinkiui priskirti atsakomybę ar kaltę DI, įtakos turi suvokiama tyčinė paslaugos nesėkmė (Sullivan & Wamba, 2022) arba suvoktas aukštesnis intelektas (Ryoo et al., 2024). Be to, nustatyta, jog suvokiamas DI neveiksnumas, gali turėti įtakos vartotojų pasitenkinimui paslaugomis (Garvey ir kt., 2023).

Rinkodaros tyrimų apžvalga rodo, kad vartotojų polinkis priskirti moralinį veiksnumą DI nėra pakankamai ištirtas, trūksta žinių, kaip šį reiškinį galėtų veikti skirtingi paslaugos subjektai – DI arba žmogus. Būtent šios priežastys identifikuoja tolimesnių tyrimų poreikį, analizuojat vartotojų polinkį priskirti moralinį veiksnumą DI arba žmogaus sprendimams paslaugose.

Ayad'as ir Plaks'as (2025) rekomenduoja plėsti mokslinių tyrimų sritį, įvertinant skirtumus tarp suvokiamo DI ir žmogaus moralinio veiksnumo identiškoje moralinėje situacijoje. Be to, autoriai pabrėžia, jog mokslinėje literatūroje analizuojamos kraštutinės moralinių situacijų aplinkybės, t.y. teigiamos arba neigiamos. Mokslininkų teigimu, trūksta tyrimų, kurie analizuotų, kaip kinta priskiriamas moralinis veiksnumas dviprasmiškose moralinėse aplinkybėse (Ayad & Plaks, 2025; Banks, 2021). Panašiai siūlo Maninger'is ir Shank'as (2022), pabrėždami, jog tolimesniuose tyrimuose svarbu nagrinėti, kaip keičiasi žmonių suvokiamas DI moralinis veiksnumas, kai priimami moraliai teisingi sprendimai. Todėl, galima teigti, jog svarbu ištirti egzistuojančius skirtumus tarp

suvokiamo DI ir žmogaus moralinio veiksnio bei kokią įtaką šiai sąlygai turi žmogui palankūs, ir nepalankūs moraliniai sprendimai.

Garvey'is ir kt.(2023) siūlo įvertinti, kaip vartotojai suvokia DI ir žmogaus sąžiningumą bei identifikuoti galimą to poveikį vartotojų ketinimams. Šiam tikslui įgyvendinti, autoriai rekomenduoja įvertinti ribines sąveikos aplinkybes – vartotojams palankias ir nepalankias. Pavone'as ir kt. (2023) rekomenduoja įvertinti galimą paslaugos teikimo nesėkmės poveikį vartotojų atsakui, pavyzdžiui pasitenkinimui. Atsižvelgiant į mokslininkų rekomendacijas, galima teigti, jog svarbu suprasti, kaip vartotojų polinkis priskirti moralinį veiksnį DI ir žmogaus sprendimams veikia pasitenkinimą paslaugomis ir kokią įtaką gali turėti sprendimo (ne)palankumas, ir suvokiamas (ne)sąžiningumas.

Mokslinių tyrimų apžvalga rinkodaros ir kitose mokslo srityse atskleidė, jog egzistuoja aiški mokslinės literatūros spraga, analizuojant vartotojų polinkį priskirti moralinį veiksnį DI. Mokslininkų apibrėžtos tolimesnių tyrimų kryptys, susijusios su moralinėje situacijoje dalyvaujančias subjektais (DI ir žmogus), moralinių sprendimų dviprasmiškumu, sprendimų palankumu ar nepalankumu, suvokiamu DI ir žmogaus sąžiningumu bei galimu poveikiu vartotojų atsakui, padėjo suformuluoti tolimesnio tyrimo poreikį. Tampa svarbu suprasti, kaip skiriasi vartotojų polinkis priskirti moralinį veiksnį DI ir žmogui moraliai abejotinosiose paslaugų aplinkybėse, kai priimtas sprendimas buvo palankus arba nepalankus vartotojo atžvilgiu. Be to, svarbu nustatyti, kokį vaidmenį atlieka suvokiamas DI arba žmogaus sąžiningumas. Galiausiai, svarbu įvertinti, kaip tai veikia vartotojų atsaką – pasitenkinimą. Toks tyrimas suteiktų naujų žinių akademinio ir praktinio požiūriu, padėsiančių įvertinti, kokiomis aplinkybėmis vartotojai yra linkę priskirti moralinį veiksnį DI ir žmogaus sprendimams bei, kaip tai veikia pasitenkinimą paslaugomis. Gauti rezultatai užtikrins efektyvesnį DI taikymą paslaugų sektoriuje.

Apibendrinant visų šioje darbo dalyje atliktų tyrimų rezultatus, identifikuojamos probleminės sritys, reikalaujančios gilesnio supratimo:

- *Kuriam paslaugų teikimo subjektui (DI ar žmogui) ir kokiomis aplinkybėmis vartotojai labiau linkę priskirti moralinį veiksnį;*
- *Kaip vartotojų polinkis priskirti moralinį veiksnį DI ir žmogaus sprendimams veikia pasitenkinimą paslaugomis.*

Atsižvelgiant į apžvelgtų tyrimų dviprasmiškumą ir ribotą kontekstualizaciją, šiame magistro projekte tyrimo objektas yra vartotojų priskiriamo moralinio veiksnio DI ir žmogaus sprendimams poveikis pasitenkinimui paslaugomis.

2. Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis teorinis pagrindimas

Šiame skyriuje analizuojami atlikti moksliniai darbai, siekiant apžvelgti moralinio veiksnumo reiškinių paslaugų kontekste, polinkį priskirti moralinį veiksnumą paaiškinančias teorijas, sprendimo palankumo ir polinkio priskirti moralinį veiksnumą sąsajas, skiriamas dėmesys suvokiamo sąžiningumo konceptualizavimui ir ryšių tarp vartotojų pasitenkinimo paslaugomis pagrindimui. Remiantis analizės rezultatais, argumentuojamas konceptualus modelis.

2.1. Moralinis veiksnumas ir jo reikšmė paslaugų kontekste

Mokslinėje literatūroje moralinis veiksnumas apibrėžiamas, kaip „[...] bet kokios rūšies subjekto (žmogaus, gyvūno, roboto ar kt.) veikimas autonomiškai, vadovaujantis normatyviniais sprendimais ir savikontrole.“ (de Pedro, 2024, p. 696). Moksliniuose darbuose išskiriamos keturios moralinio veiksnumo sąlygos, kurios nusako moralinio agento (angl. *moral agent*) gebėjimus būti racionaliam, veikti autonomiškai, turėti laisvą valią, t.y. elgtis nepriklausomai laisvai, ir fenomenalią sąmonę, t.y. pasižymėti subjektyvia patirtimi (Zafar, 2024). Panašiai teigia ir Firt'as (2024), pažymėdamas, jog subjekto moralinis veiksnumas nusako moralinio agento sąmoningumą, laisvą valią, gebėjimą moraliai mąstyti, suprasti bei svarstyti, taikant teisingus moralės principus. Moralinis agentas turi pasižymėti panašiomis į žmogaus psichinėmis savybėmis, kurios leistų suprasti bei atskirti priimamų moralinių sprendimų teigiamą ir / ar neigiamą poveikį. „Suvoktas veiksnumas kvalifikuoja subjektus kaip moralinius agentus, kurie gali daryti gera arba bloga.“ (Sullivan & Wamba, 2022). Apibendrinant moralinio veiksnumo apibrėžimus, galima teigti, jog paprastai, moralinis veiksnumas apima subjekto gebėjimą veikti laisvai, savarankiškai, tikslingai bei moraliai.

Egzistuoja skirtingi požiūriai į dirbtinio intelekto technologijomis paremtus moralinius subjektus. Galima teigti, jog DI kvestionuoja tradicinės moralinio veiksnumo sampratą, nes, anot mokslininkų, „[...] žmonės nebėra vieninteliai subjektai, galintys pasižymėti tokiu protu ir racionalumu, kad galėtų sąmoningai veikti moraliai.“ (Llorca Albareda, 2024, p. 2). Dalis mokslinės bendruomenės narių teigia, kad DI moralinis veiksnumas reikšmingai priklauso nuo technologijos savarankiškumo lygio. Tarpdisciplininame informatikos, psichologijos, etikos ir filosofijos mokslų lauke, dirbtinių moralinių agentų (DMA) (angl. *artificial moral agent*, AMA) sąvoka apibrėžiama, kaip bet kokio pobūdžio DI technologija (Zafar, 2024), galinti priimti moralinius sprendimus, atitinkančius turimo autonomiškumo lygį (Vijayaraghavan & Badea, 2024). Panašiai teigia Sullivan ir Wamba (2022), kai DI technologija suvokiama, kaip gebanti jausti patirtines psichines būsenas ar turėti fenomenalią sąmonę, ji gali būti pripažinta dirbtiniu moraliniu agentu, kuris gali įvertinti situaciją remiantis moraliniais pagrindais (Sullivan & Wamba, 2022). Vijayaraghavan'o ir Badea (2024) teigimu, egzistuoja trys DMA kategorijos, kurios nusakomos įdiegtų moralinių principų ir autonomiškumo lygmeniu:

1. Numanomi (angl. *implicit*) DMA. DI technologijos konstrukcijoje įdiegti moraliniai pagrindai, tačiau sistemos nėra pajėgios vertinti moralinės situacijos teigiamai arba neigiamai;
2. Apibrėžti (angl. *explicit*) DMA. DI sistemose algoritmiškai integruoti moraliniai principai ir etinės taisyklės, tiksliai apribojančios ir nurodančios moralinės elgsenos kryptį;
3. Visiški (angl. *full*) DMA. DI, pasižymintys žmogui artimomis savybėmis, dideliu autonomiškumo lygiu, pagrįstais veiklos motyvais bei laisva valia (Vijayaraghavan & Badea, 2024).

Tuo tarpu Formosa'as ir Ryan'as (2021), remdamiesi mokslininko Moor'o (2009) keturių tipų moralinių robotų koncepcija, kvalifikuoja pirmojo, antrojo, trečiojo ir ketvirtojo lygio robotus (angl. *bots*). Anot autorių, terminas robotas, apima visas DMA kategorijas, turinčias fizinį ir virtualų pavidalą. Šiuo atveju botams priskirti lygiai, apibrėžia jų kategoriją, savarankiškumą, interaktyvumą, gebėjimą moraliai elgtis ir mąstyti:

1. Pirmojo lygio DMA. Šie robotai apima beveik bet kokio pavidalo dirbtinius subjektus, kurie gali daryti etinį poveikį, t.y. sukelti moralines pasekmes. Tačiau, neturi arba pasižymi itin mažu autonomiškumu, interaktyvumu, todėl negali prisitaikyti ir atlikti inicijuotų sprendimų situacijai spręsti;
2. Antrojo lygio DMA. Šie robotai pasižymi etinėmis savybėmis, o daugiausiai saugumu ir patikimumu, t.y. veikia pagal iš anksto įdiegtas etines sąlygas. Tačiau, minėtomis sąlygomis nesivadovauja, siekiant sąžiningumo ir negali reaguoti į kitokias, nei numatyta, moralines aplinkybes;
3. Trečiojo lygio DMA. Autorių teigimu, trečiojo lygio robotai „[...] turi aiškia, etinių taisyklių, normų ar dorybių reprezentaciją ir gali nuspręsti arba apskaičiuoti, ką tame kontekste yra morališkai gera daryti, ir veikti remdamiesi šiuo moraliniu sprendimu.“ (Formosa & Ryan, 2021, p. 840). Nepaisant to, 3-io lygio DMA, vertindamas tam tikrą moralinę situaciją, gali ją spręsti platesne galimų variantų įvairove, apimant labai bendrinius ir itin specifinius;
4. Ketvirtojo lygio DMA. Ketvirtojo lygio robotai turi sąmonę ir pasižymi tokiomis pat kognityvinėmis savybėmis, kaip žmogus. Pasižymi gebėjimu savarankiškai priimti sprendimus, neveikiamas jokių išorinių jėgų, atskirti gėrį nuo blogio ir elgtis moraliai visuose kontekstuose (Formosa & Ryan, 2021).

Taigi, dažniausiai, DMA moralinis veiksnumas priklauso nuo gebėjimo veikti moraliai, tikslingai bei turimo autonomiškumo lygio. DI autonomiškumo svarbą, vertinant jį kaip aukščiausio lygmens dirbtinį moralinį agentą, pabrėžia ir Firt'as (2024), pažymėdamas, kad būtent ši technologijų kategorija, geba savarankiškai atlikti „[...] sudėtingas užduotis, įskaitant reikalaujančias žmogaus lygio kognityvinių gebėjimų; turi moralinį supratimą, t. y. geba veikti remdamiesi morale, o ne tik pagal moralę...“ (Firt, 2024, p. 176). Panašiai teigia Zafar'as (2024), autonomiškumo svarbą papildydamas dar dvejomis DMA savybėmis, kurios, anot jo, yra itin svarbios, vertinant DI, kaip moralinį agentą. Šios savybės apima gebėjimą teikti prasmingus rezultatus ir reikšmingas rekomendacijas sprendimų priėmimo procesuose. Vis dėlto, mokslininkai pabrėžia, kad DI moralinis veiksnumas turėtų būti vertinamas, kaip kontinuumas, nes ateityje atsiras „įvairių sistemų, pasižyminčių skirtingais [moralinio] išprusimo lygiais.“ (Formosa & Ryan, 2021, p. 841, cit. iš Asaro, 2006, p. 11).

Pastebima, jog akademinėje bendruomenėje eksponentiškai auga susidomėjimas DMA vertinimu (Harris & Anthis, 2021), tačiau mokslinių tyrimų išvados yra dviprasmiškos. Viena vertus, dalis autorių teigia, jog DMA gali pasižymėti moraliniu veiksnumu ir jis gali būti vertinamas kaip lygiavertis moralinis agentas. Coeckelbergh'as (2016) savo moksliniame darbe patvirtino, kad autonominiai savaeigiai automobiliai gali būti vertinami, kaip moraliniai agentai, moralinį veiksnumą visapusiškai priskiriant DI paremtai technologijai. Be to, Gamez ir kt. (2020), remdamiesi dorybių etika, kaip moralės teorija, teigia, kad DI gali būti moralus ir elgtis etiškai, todėl jis gali būti vertinamas, kaip lygiavertis arba panašus į gyvuosius moralinius agentus. Kita vertus, dalis mokslininkų teigia priešingai, pabrėždami, jog DMA yra „[...] hipotetinis neįsikūnijusių, labai sudėtingų, giliuoju mokymusi grindžiamų dirbtinio intelekto asistentų...“, kurių perspektyvos

tolimoje ateityje yra galimos, tačiau iki šiol nerealizuotos (Constantinescu et al., 2022, p. 2). Teigiama, kad bent kol kas, DI jokiais sąlygomis negali būti vertinamas kaip moralinis agentas, nes jis neatitinka pagrindinių moralinio veiksnio sąlygų. Visa tai grindžiama Constantinescu ir kt. (2022) tyrimo išvadomis, kurios nurodo, jog DMA negali ir dar kurį laiką negalės būti atsakingais moraliniais agentais, o moralinė atsakomybė yra viena iš moralinio veiksnio dedamųjų. Taip pat, mokslininkai Sison'as ir Redín (2021), remdamiesi neoaristoteliška perspektyva, atskleidė, jog DI negali būti laikomi moraliniais agentais, dėl negebėjimo veikti autonomiškai bei priimti moralinių sprendimų. Teigiama, jog moralinis teisingumas yra vidinė savybė, kuri nėra įmanoma DI atžvilgiu (Sison & Redín, 2023). Panašiai teigia Véliz (2021), jog DI negali jausti ir atlikti moralinių vertinimų, tačiau jie gali elgtis, kaip moraliniai agentai, todėl būtent tokios būtybės apibūdinamos, kaip „nenuoseklūs moraliniai subjektai“.

Nors mokslinėje bendruomenėje egzistuoja prieštaringi požiūriai, vertinant DI, kaip moraliniu veiksnio pasižyminčius moralinius agentus, autoriai pabrėžia, kad žmonės, socialinėse sąveikose dalyvaujančias technologijas, yra linkę moraliai vertinti (Harris & Anthis, 2021, Tavani, 2018). Žmonių polinkis priskirti moralinį veiksnį socialiniams robotams, įskaitant, bet neapsiribojant DI, apibrėžiamas moralinio veiksnio suvokimo (SMV) (angl. *perceived moral agency*, PMA) sąvoka (Banks et al., 2023). „Moralinis veiksnys grindžiamas suvokimu, kad dirbtinis agentas turi moralinį gebėjimą (t. y. įgimtą gerumo ar blogumo potencialą) ir veiksnio gebėjimą (t. y. gebėjimą veikti moraliai savo nuožiūra, nepriklausomai nuo jo kūrėjo)“ (Banks et al., 2023, p. 32:4, cit. iš Banks, 2019). Todėl, galima patvirtinti, jog nepaisant to, ar DI grįsta technologija išties gali pasižymėti moraliniu veiksnio, žmonės yra linkę priskirti tam tikro lygio moralinį statusą. Tampa svarbu suprasti, kas ir kaip lemia žmonių polinkį priskirti moralinį veiksnį DI ir žmogaus sprendimams.

Mokslininkai nagrinėja atskiras SMV dedamąsias dirbtinių agentų atžvilgiu. Anot Banks (2019), SMV negali egzistuoti be būtinų suvokiamo moralumo ir veiksnio sąlygų, todėl svarbu suprasti, kas turi įtakos žmonių polinkiui priskirti moralę arba veiksnį, įvertinant šiuos veiksnis, kaip atskirus vienetus. Suvokiamas dirbtinių agentų moralumas daugiausiai analizuotas psichologijos, sociologijos, etikos, žmogaus ir kompiuterio sąveikos mokslo srityse. Naujausi moksliniai tyrimai nagrinėja, kaip kinta žmonių suvokiamas moralumas, sąveikaujant su dirbtiniais agentais ir žmonėmis identiškose moralinėse aplinkybėse. Bigman'o ir kt. (2022) tyrimu siekta išsiaiškinti, kaip žmonės vertina algoritmo ir žmogaus priimtus moralinius sprendimus, kurie gali diskriminuoti tam tikras asmenų grupes. Rezultatai parodė, kad žmonės skirtingai suvokia algoritmo ir žmogaus moralumą diskriminacijos atveju. Žmonės palankiau reaguoja į algoritmo priimtus diskriminacinio pobūdžio sprendimus. Tam įtakos turi suvokiamos DI ir žmogaus išankstinės nuostatos bei tyčiniai ar netyčiniai ketinimai diskriminuoti, t.y. žmonės yra linkę manyti, kad algoritmas negali inicijuotai pažeisti moralės normų, tuo tarpu žmogus – gali.

Wilson ir kt. (2022) nagrinėjo, kaip žmonės suvokia DI ir žmogaus moralumą, moraliai neteisingose situacijose. Tyrimas parodė, kad žmonių ir DI priimti amoralūs sprendimai buvo suvokiami nevienodai – žmonės buvo vertinami griežčiau nei DI. Žmonių sprendimai buvo suvokiami kaip labiau neteisingi nei DI, be to, DI amoralūs sprendimai suvokiami, kaip labiau leistini nei žmogaus.

Panašiai teigia ir Shank'o (2019) tyrimo išvados, kurios nurodo, kad žmonės atlaidžiau vertina DI moralinius nusižengimus, lyginant su identiškais žmogaus moraliniais sprendimais. Nors DI

suvokiami, kaip mažiau moraliai atsakingi nei žmonės, jų sprendimai yra vertinami moraliniu požiūriu ir turi įtakos tolimesnei žmonių reakcijai.

Tuo tarpu rinkodaros ir susijusiose mokslo srityse, pastebimas augantis mokslininkų susidomėjimas suvokiamo dirbtinių agentų moralumo tema, tačiau tyrimų išvados yra nevienareikšmės. Söderlund'as (2023b) analizavo kokią įtaką suvokiamam DI paslaugų roboto moralumui, žmogiškumui, vertinimui ir reakcijai bendrąja prasme, turi privatumo pažeidimai. Rezultatai atskleidė, jog žmonės yra linkę suvokti robotus, kaip moraliai veikiančius subjektus. Amoralūs roboto veiksmai turėjo įtakos žmonių suvokiamam mažesniai roboto žmogiškumui (antropomorfizmui). Suvokiamas roboto moralumas ir žmogiškumas teigiamai veikia vartotojų pasitikėjimą ir robotų priimtinumą paslaugose. Be to, tyrimas atskleidė, kad paslaugų roboto veiksmai vertinami, remiantis tokiais pat moraliniais principais, kaip ir žmogaus. Tuo tarpu Gill'o (2020) tyrimas parodė, kad vartotojai taiko skirtingas moralines normas automobilius vairuojantiems žmonėms ir autonominėms transporto priemonėms.

Dietvorst'as ir Bartels'as (2022) nagrinėjo vartotojų suvokiamą algoritmo ir jo sprendimų priimtinumą moralinių aplinkybių kontekste. Tyrimo rezultatai atskleidė, jog vartotojai neigiamai vertina algoritmų sprendimus moralinėse situacijose. Visa tai lemia vartotojų požiūris, jog algoritmai taiko sprendimų maksimizavimo strategiją, kuomet galutinio tikslo (pelno, efektyvumo ar kt.) siekiama neatsižvelgiant į kitus, moraliai svarbius veiksnius (teisingumą, sąžiningumą ar kt.). Vartotojų nuomone, ši strategija taikoma ir paslaugą teikiančių žmonių tarpe, tačiau lyginant su algoritmu, jo sprendimai vertinami griežčiau. Todėl, moraliniu požiūriu svarbiose paslaugos teikimo aplinkybėse, vartotojai pageidauja, kad sprendimą priimtų žmogus, nes algoritmai suvokiami, kaip negalintys priimti moraliai teisingų sprendimų.

Tuo tarpu suvokiamas dirbtinių agentų veiksnumas nagrinėjamas naujausiuose tarpdisciplininių mokslo krypčių tyrimuose. Geiselmann ir kt. (2023) analizavo žmonių polinkį priskirti žmogiškas savybes DI, įskaitant veiksnumą. Tyrimas parodė, kad žmonių manymų, DI gali pasižymėti veiksnumu, t.y. gebėjimu veikti savo nuožiūra, turėti ketinimų ir priimti inicijuotus sprendimus. Polinkiui priskirti veiksnumą, įtakos turi žmogaus pavidalui artima dirbtinio agento išvaizda. Be to, suvokiamas veiksnumas veikiamas kontekstinių aplinkybių bei atliekamų užduočių pobūdžio. Atliekant užduotis, kurios paprastai atliekamos žmonių, pavyzdžiui priimant sprendimus socialinėje sąveikoje, didina suvokiamą dirbtinio agento veiksnumą. Tuo tarpu mechaninės užduotys veikia priešingai – mažina šią prielaidą.

Swiderska'os ir Küster'io (2020) moksliniu darbu siekta atskleisti, kaip žmonės suvokia žalingai besielgiančius robotus moraliniu požiūriu, įvertinant polinkį priskirti veiksnumą. Tyrimas parodė, jog žmonės buvo linkę priskirti didesnį veiksnumą robotams, kurie demonstravo geranorišką elgseną nei tiems, kurių veiksmai turėjo žalingų moralinių pasekmių. Ši išvada grindžia žmonių polinkį priskirti veiksnumą robotams, kurių veiksmai yra teigiami žmonių atžvilgiu. Be to, nustatyta, kad suvokiamas roboto veiksnumas turi teigiamos įtakos pasitikėjimui. Lyginant suvokiamą roboto veiksnumą su žmogumi, didesnis veiksnumas bet koku atveju buvo priskiriamas žmogui.

Zafari ir Koeszegi (2021) nagrinėjo, kaip suvoktas roboto veiksnumas veikia žmonių reakciją ir požiūrį bei kokią įtaką šiam reiškiniui turi suvokta kontrolė. Nustatyta, kad suvokiamas roboto gebėjimas savarankiškai priimti sprendimus, neigiamai veikia žmonių reakciją į robotą – identifiktuotas mažesnis pasitikėjimas ir didesnė numanomos grėsmės tikimybė. Mažesnis

pasitikėjimas lemiamas suvoktu roboto nenuspėjamumu ir gebėjimų veikti nepriklausomai nuo žmonių. Suvokiamai grėsmei įtakos turėjo kontrolės stoka, t.y. negebėjimas kontroliuoti roboto sprendimų.

Suvokiamas veiksnumas nagrinėjamas ir rinkodaros bei susijusiose mokslo srityse. Pitardi ir kt. (2022) moksliniame darbe analizuota suvokiamo DI paslaugų roboto veiksnumo įtaka vartotojų pasitikėjimui, naudojantis gėdingomis (angl. *embarrassing*) paslaugomis, pavyzdžiui, konsultuojantis dėl lytiškai plintančių ligų. Tyrimas parodė, kad vartotojai priskiria mažesnę veiksnumą DI nei žmonėms. Būtent tokiose paslaugos teikimo aplinkybėse, suvoktas paslaugų roboto neveiksnumas teigiamai veikia vartotojų reakciją ir DI priimtinumą, lyginant su žmonėmis. Vartotojų nuomone, paslaugą teikiantys žmonės gali susidaryti neigiamą nuomonę ir tai didina jaučiamą gėdos jausmą. Tuo tarpu DI suvokiamas, kaip mechaninis įrankis, kuris negali vertinti situacijos, todėl gėdos jausmas sumažėja.

Garvey'is ir kt. (2023) savo tyrimu analizavo, kaip vartotojai vertina DI ir žmogaus pateiktus rinkodaros pasiūlymus, įvertindami suvokiamą subjekto veiksnumą. Tyrimo rezultatai atskleidė, jog vartotojai nėra linkę priskirti veiksnumo DI, nes, anot jų, dirbtiniai agentai negali suvokti bei patys nuspręsti dėl savo sprendimų priimtumo kliento atžvilgiu. Todėl, vartotojų manymu, DI pateikti rinkodaros sprendimai, lyginant su žmogaus pateiktais identiškais pasiūlymais, yra moralesni, nesavanaudiški.

Naujausi moksliniai tyrimai patvirtina, kad suvoktas subjekto moralumas ir veiksnumas, yra reikšmingai susiję (Trafton et al., 2023). Trafton'o ir kt. (2024) teigimu, svarbu toliau nagrinėti suvokiamą moralinį veiksnumą, kaip vientisą reiškinį, ypatingą dėmesį skiriant žmogaus ir dirbtinių agentų sąveikai. Dar 2019 metais, mokslininkė Banks sukūrė dviejų dimensijų skalę, skirtą vertinti žmonių priskiriamą SMV bet kokios kategorijos dirbtiniams agentams ir žmonėms. Pirmoji suvokiamo moralinio veiksnumo dimensija „moralė“, nurodo žmogaus polinkį priskirti moralines ir etines savybes socialinėje sąveikoje dalyvaujantiems dirbtiniams agentams ir žmonėms. Suvokiamai moralei įvertinti, išskirti šeši elementai, kuriais matuojama subjekto nuovoka, vertinant moralinę situaciją gėrio ir blogio aspektais, gebėjimas apgalvoti priimto sprendimo moralumą, racionaliai vertinti moralinę situaciją, elgtis, remiantis moralinėmis normomis, susilaikyti nuo moralinių pasekmių turinčių veiksmų bei išsipareigoti moraliai elgsenai (Banks, 2019). Skalės validumui skirtų mokslinių tyrimų metu, antroji skalės dimensija kito iš suvokiamo veiksnumo į suvokiamą priklausomybę. Autorė pabrėžia, jog formuojant skalės dimensijas ir įvertinus veiksnumą, kaip galimą skalės konstrukta, buvo nustatyta, kad žmonės rečiau vertina dirbtinių agentų savarankiškumą ir dažniau priklausomybę nuo aplinkos. Todėl, dimensija „priklausomybė“, matuoja suvokiamą subjekto – dirbtinio agento ir žmogaus – priklausomybę nuo išorinių veiksmų, pavyzdžiui programavimo ar kitų subjektų. Šiai skalės dimensijai matuoti išskirti keturi elementai, kuriais vertinamas suvokiamas subjekto gebėjimas elgtis tik taip, kaip yra užprogramuotas, daryti tai, ką įpareigoja atlikti žmonės, atsisakyti daryti tai, kas užprogramuota bei suvokiamas atlikto veiksmo rezultatas, kurį nulėmė programavimas (Banks, 2019). Autorės Banks (2019) sudarytoje lentelėje (žr. **1 priede**), nurodomi SMV veiksnus reprezentuojantys elementai – (1) moralė, (2) priklausomybė. Lentelėje esanti [x] reikšmė yra kintanti, atsižvelgiant į agento tipą, t.y. robotą, algoritmą, kitos kategorijos technologiją, ar žmogų.

SMV plačiausiai nagrinėjamas naujausiuose žmogaus ir kompiuterio sąveikos tyrimuose. Wester'io ir kt. (2024) tyrimu siekta įvertinti SMV įtaką psichikos sveikatos pacientų pasitikėjimui,

suvokiamam saugumui ir pokalbių roboto priimtinumui. Autorių teigimu, suvokiamas dirbtiniu agentų moralinis veiksnumas, tokiose jautriose paslaugose, kaip psichikos sveikata, yra itin svarbus veiksnys, padedantis suprasti, kokios dirbtinio agento savybės gali padidinti pasitikėjimą, užtikrinti saugumą ir paskatinti naudojimąsi pokalbių robotais medicinos srityje. Eksperimentinio tyrimo rezultatai atskleidė, kad žmonės yra linkę priimtiniau reaguoti, labiau pasitikėti ir jaustis saugiau, komunikuojant su pokalbių robotais, kurie pasižymi didesniu SMV. Suvokiamam moraliniam veiksnui įtakos turi roboto panašumas į žmogų (antropomorfizmas), gebėjimas įsijausti į situaciją, būti empatišku ir dėmesingu (šiluma), supratingu, ir švelniu (jautrumas). Be to, analizuojant šį reiškinį paauglių tarpe, išskirtos papildomos savybės, prisidedančios prie suvokiamo didesnio pokalbių roboto moralinio veiksnio. Minima, kad pokalbių robotas, komunikodamas su paaugliais, turėtų gebėti prisitaikyti prie jaunuolių vartojamo žargono, gebėti išreikšti humorą ir išlaikyti draugišką, ramų toną, kuris padėtų sukurti artimesnį tarpusavio ryšį. Tuo tarpu suvokiamam mažesniai moraliniam veiksnui įtakos turi pernelyg tiesmukiškas, klinikiniais terminais paremtas bendravimas (Wester et al., 2024). Autoriai pabrėžia, kad tolimesni tyrimai turėtų nagrinėti, kaip SMV skiriasi sąveikaujant su pažangesniais dirbtiniais agentais, kurie galėtų pateikti natūralesnius atsakymus psichikos sveikatos klausimais. Taip pat, įvertinti ilgalaikį SMV pokalbių robotų poveikį žmonių pasitikėjimui ir ketinimams naudotis šia technologija psichikos sveikatos paslaugose.

Nijssen ir kt. (2022) savo tyrimu nagrinėjo sąsajas tarp suvokiamo roboto veiksnio, afektinių būsenų ir moralumo, siekiant įvertinti žmonių polinkį pasitikėti roboto priimtais moraliniais sprendimais. Tyrimas parodė, jog visi minėti veiksniai reikšmingai veikia žmonių pasitikėjimą moralinių sprendimų priėmimo kontekste, t.y. didėjant suvokiamam roboto veiksnui ar gebėjimui jausti patirtines emocijas, auga žmonių pasitikėjimas. Teigiama, jog šis ryšys yra veikiamas psichologinių veiksnių, kuomet suvokiamos roboto antropomorfizmo savybės teigiamai veikia požiūrį į robotą, vertinant jį, kaip moraliai atsakingą subjektą. Tačiau, nustatyta, jog abiejų savybių, antropomorfizmo ir veiksnio derinys nekelia didesnio žmonių pasitikėjimo. Todėl, daroma prielaida, jog per didelis panašumas į žmogų, gali kelti neigiamą žmonių reakciją, pavyzdžiui nerimą. Remiantis tyrimo išvadomis, galima teigti, kad suvokiamam dirbtinio agento moraliniam veiksnui reikšmingą įtaką daro ir afektinių būsenų suteikimas – gebėjimas jausti emocijas.

Naujausi moksliniai tyrimai nagrinėja žmonių polinkį priskirti moralinį veiksnumą dirbtiniam agentui ir žmogui. Manoli ir kt. (2024) analizavo žmonių polinkį priskirti teigiamą arba neigiamą moralinį veiksnumą DI ir žmogui, moraliai svarbių aplinkybių kontekste. Nustatyta, kad žmonės yra linkę suvokti DI pokalbių robotus, kaip moraliniu veiksniumu pasižyminčius subjektus, t.y. gebančius priimti moraliai teisingus arba neteisingus sprendimus. Tarpusavyje lyginant DI ir žmogaus sprendimų (a)moralumą, pastebėta, jog subjektams taikomi dvigubi standartai. Tyrimas atskleidė, jog žmonės griežčiau vertina DI priimtus moralinius sprendimus ir yra linkę manyti, kad vienam DI agentui nusižengus, visi kiti DI agentai taip pat bus linkę elgtis moraliai neteisingai. Tuo tarpu vertinant žmogaus moralinius sprendimus, nustatyta, kad vienam žmogui nusižengus, visi kiti žmonės nebus vertinami atitinkamai taip pat. Todėl, galima teigti, jog žmonės skirtingai vertina DI ir žmogaus sprendimus, generalizuodami DI amoralumą bei vertindami priimtus sprendimus griežčiau. Būtent amoralios aplinkybės veikia žmonių pasitikėjimą DI sprendimais plačiąja prasme, t.y. vienam DI sprendimui nusižengus, nepasitikima visomis kitomis DI technologijomis. Tačiau, tyrimas taip pat parodė, kad vertinant moraliai svarbiuose situacijose dalyvaujančius DI ir žmones bendrąja prasme, manoma, jog DI elgsis teisingiau nei žmogus. Visa tai lemia žmonių suvokiamas DI sąžiningumas,

nes, anot tyrimo dalyvių, žmonės, moralinėse situacijose gali elgtis savanaudiškai (Manoli et al., 2024).

Polinkio priskirti moralumą tyrimų apžvalgą atskleidė, jog žmonės yra linkę suvokti dirbtinius agentus, kaip mažiau moraliai atsakingus subjektus. Be to, dirbtiniai agentai ir jų sprendimai moraliniu požiūriu vertinami priimtinau. Tačiau, vis dar nėra aišku, kokią įtaką suvokiamas moralumas daro vartotojų ir dirbtinių agentų sąveikoje. Vieni autoriai pabrėžia, kad vartotojai yra linkę taikyti tuos pačius moralinius principus, kaip žmonėms (Söderlund, 2023b), kiti teigia priešingai (Gill, 2020). Be to, trūksta žinių apie tai, kokiomis aplinkybėmis vartotojai suvokia dirbtinius agentus, kaip moraliai veikiančius subjektus ir kada suvokiamo moralumo sąlyga atmetama. Polinkio priskirti veiksnumą dirbtiniams agentams tyrimų apžvalgą parodė, jog šiuo požiūriu egzistuoja nevienareikšmiai mokslinių tyrimų rezultatai. Vieni autoriai teigia, kad žmonės gali priskirti veiksnumą dirbtiniams agentams (Zafari & Koeszegi, 2023; Swiderska & Küster, 2021), kiti teigia priešingai (Pitardi et al., 2022; Garvey et al., 2023). Tačiau, autoriai vieningai sutaria, jog suvoktas subjekto (ne)veiksnumas turi įtakos žmonių reakcijai (Garvey et al., 2023), pasitikėjimui (Swiderska & Küster, 2020; Koeszegi, 2021) ir paslaugos priimtinumui (Pitardi et al., 2022).

Apibendrinant galima teigti, jog vartotojai yra linkę priskirti moralinį veiksnumą ne tik žmonėms, bet ir dirbtiniams agentams. Tyrimų išvados rodo, jog žmonės remiasi tais pačiais psichologiniais mechanizmais, kaip ir biologinių būtybių atveju, suvokdami dirbtinių agentų, įskaitant DI, moralinį veiksnumą (Ladak et al., 2024). Nepaisant to, žmonės gali skirtingai suvokti ir priskirti moralinį veiksnumą DI ir žmogui. Tai lemia kontekstinių aplinkybių skirtumai, sąveikos pobūdis, žmonių moraliniai principai bei asmeninės savybės. Tačiau, mokslininkai sutaria, jog polinkis priskirti moralinį veiksnumą, turi įtakos žmonių reakcijai. Autoriai pabrėžia, kad šis reiškinys stipriai koreliuoja su teigiamomis ir neigiamomis emocijomis, pavyzdžiui didesniu / mažesniu pasitikėjimu, didesniu / mažesniu saugumu, didesniu / mažesniu nerimu ir t.t.. Nepaisant akivaizdaus poveikio žmonių reakcijai, moralinio veiksnumo reiškinys mažai nagrinėtas rinkodaros mokslo srityje, ypač vartotojo ir DI sąveikoje. Naujausi tyrimai rodo, kad suvoktas subjekto (ne)veiksnumas ir (ne)moralumas gali turėti įtakos vartotojų pasitikėjimui, paslaugos priimtinumui. Autoriai rekomenduoja plėsti šio pobūdžio tyrimų kryptį, įvertinant, kaip polinkis priskirti moralinį veiksnumą kinta identiškoje vartotojo ir žmogaus sąveikoje. Todėl, tampa svarbu suprasti, kaip vartotojų polinkis priskirti moralinį veiksnumą DI ir žmogui veikia vartotojų atsaką paslaugose, nagrinėjant dar menkai tirtą poveikį pasitenkinimui.

2.2. Polinkį priskirti moralinį veiksnumą DI ir žmogaus sprendimams paaiškinančios teorijos

Suvokiamo moralinio veiksnumo reiškinys gali būti siejamas su daugeliu susijusių veiksnių, kaip žmogaus asmenybė ar kontekstas (Banks, 2020) ir gali būti priskiriamas ne tik žmonėms, bet, ir dirbtiniams agentams, pavyzdžiui DI (Banks, 2021). Moralinio veiksnumo būtiniosios sąlygos apima abi šio reiškinio dedamąsias, t.y. suvokiamą subjekto veiksnumą ir moralumą. Polinkis priskirti veiksnumą siejamas su 1978 m. Premack'o ir Woodruff'o pasiūlyta proto teorija (PT) (angl. *theory of mind*, TOM). Ši teorija nurodo, kad žmonės yra linkę priskirti psichines būsenas sau ir kitiems, jas lyginti tarpusavyje, siekiant nuspėti galimą stebimo subjekto elgesį, turimas žinias, įsitikinimus ar ketinimus (Banks, 2020).

Mokslinėje literatūroje atskleidžiama, jog polinkis priskirti psichinės būsenas yra grindžiamas žmogaus suvokimo procesais, todėl moksliniuose darbuose *proto suvokimas* ir *proto priskyrimas*

naudojami, kaip sinonimai (Galvez & Hanono, 2024). Pabrėžiama, jog šių sąvokų prasmė teoriniu požiūriu nesiskiria (Thellman et al., 2022). Todėl, šiame projekte šie terminai bus naudojami pakaitomis.

Teigiama, jog suvokiamas protas lemia žmonių polinkį priskirti proto gebėjimus stebimam subjektui (Esterwood & Robert, 2023). Akcentuojama, jog šis procesas siejamas su egocentrišku šališkumu, kuomet žmonės yra linkę suvokti kitų proto gebėjimus, remiantis savo paties proto modeliu (Söderlund, 2022). Anot Banks (2020), suvokiamo proto reiškinyje yra nesąmoningas, automatiškas ir spontaniškas, lemiamas kognityvinių ir socialinių reakcijų.

Suvokiamas protas yra glaudžiai siejamas ne tik su suvokiamu subjekto veiksmu, bet ir patirtimi. Pasak Esterwood'o ir Robert'o (2023), žmonės yra linkę nesąmoningai skirstyti suvokiamą protą į dvi kategorijas, t.y. suvokiamą patirtį ir suvokiamą veiksmą. Suvokiama patirtis siejama su subjekto gebėjimu jausti emocijas, pavyzdžiui liūdesį ar džiaugsmą. Tuo tarpu suvokiamas veiksmas apibrėžia subjekto gebėjimą sąmoningai mąstyti, priimti sprendimus, veikti turint ketinimų bei kontroliuoti savo veiksmus (Shank et al., 2021).

Pabrėžiama, jog abi proto suvokimo kategorijos gali būti priskiriamos nepriklausomai viena nuo kitos arba kartu (Esterwood & Robert, 2023). Teigiama, kad žmonės yra linkę priskirti protą ir tiems, kurie iš esmės negali pasižymėti šia biologine savybe. Banks (2020) atliktas tyrimas parodė, kad žmonės yra linkę nesąmoningai suvokti socialinėje sąveikoje dalyvaujančius dirbtinius agentus, kaip galinčius pasižymėti panašiomis psichinėmis savybėmis į žmogų. Be to, šį reiškinį galima paskatinti formuojant žmonių įsitikinimus apie dirbtinio agento proto galimybes, manipuliuojant dirbtinio agento išvaizda ar elgsena (Li et al., 2022). Mokslinėje literatūroje suvokiamo proto reiškinyje nagrinėtas analizuojant žmonių polinkį priskirti šias savybes skirtingo tipo dirbtiniams agentams, pavyzdžiui robotams (Müller et al., 2021; Yam et al., 2021, Li et al., 2022) ar DI (Clegg et al., 2024; Li et al., 2023; Lee & Hahn, 2024; Stein et al., 2020). Vis dėlto, ankstesnių tyrimų rezultatai patvirtina, jog lyginant žmonių polinkį priskirti psichikos būsenas žmonėms ir dirbtiniams agentams, didžioji dalis dažniau šias savybes priskiria žmonėms (Thellman et al., 2022). Lyginant suvokiamą dirbtinių agentų patirtį ir veiksmą, žmonės yra linkę manyti, kad pastarieji pasižymi didesniu veiksmu nei patirtimi, todėl gali smerkti moraliai neteisingo elgesio atveju (Lee et al., 2021).

2007 m., psichologijos ir neuromokslo tyrėjai Gray ir kt., sukūrė dviejų dimensijų skalę, skirtą matuoti suvokiamą žmonių ir kitų būtybių, įskaitant dirbtinių agentų, protą (žr. **2 priede**). Pirmoji suvokiamos patirties dimensija matuojama remiantis vienuolika veiksnių. Analizuojamas suvokiamas bado jausmas, gebėjimas jausti baimę, skausmą, malonumą, pyktį, troškimą, pasididžiavimą, gėdą, džiaugsmą, pasižymėti sąmone bei turėti asmenybę. Antroji dimensija skirta matuoti suvokiamą subjekto veiksmą. Šiam reiškiniui matuoti, remiamasi septyniais veiksniais. Nagrinėjama suvokiama savikontrolė, moralė, atmintis, gebėjimas atpažinti emocijas, planuoti savo veiksmus, bendrauti ir mąstyti.

PT yra viena esminių polinkio priskirti veiksmą tyrimų dalis. Naujausiuose moksliniuose darbuose ji plačiausiai nagrinėjama žmogaus ir kompiuterio sąveikos, psichologijos, žmonių elgsenos tyrimuose. Pavyzdžiui, Hwang ir Won (2022), remiantis PT sistema, savo darbe nagrinėjo suvokiamą veiksmą ir jo poveikį žmogaus, ir skirtingų dirbtinių agentų tipų sąveikoje. Autorės, atsižvelgdamos į PT, išskyrė dvi susijusias kategorijas – patirtį ir veiksmą – kaip žmonių suvokiamo proto dimensijas. Tyrimas atskleidė, jog žmonės pasižymi didžiausia suvokiama patirtimi

ir veiksnumu, tuo tarpu skirtingi dirbtinių agentai šiuo požiūriu yra vertinami nevienareikšmiškai. Šiek tiek mažesnis nei žmonėms ir didžiausias veiksnumas skirtingų dirbtinių agentų tarpe priskiriamas DI. Vėliau, polinkis priskirti veiksnumo savybes mažėjimo tvarka išsidėstė taip – robotams, virtualiems asistentams („Apple Siri“, „Amazon Alexa“, „Google Assistant“) ir galiausiai kompiuteriams. Be to, didžiausias suvokiamas veiksnumas buvo priskirtas tiems DI, kurie yra didesnės technologinės ekosistemos dalis (Hwang & Won, 2022). Tačiau, tyrėjos pabrėžia, jog pernelyg didelis suvokiamas veiksnumas ir per maža suvokiama patirtis gali turėti neigiamą poveikį žmonių pasitikėjimui. Priešingu atveju, kyla rizika pernelyg dideliame pasitikėjime ir DI galimybių pervertinimui.

Di Dio ir kt. (2020), remdamiesi PT, tyrė, kaip suvokiamas veiksnumas veikia vaikų pasitikėjimą robotu, lyginant identiškas sąveikos sąlygas su žmogumi. Tyrimo rezultatai parodė, jog vaikai vienodai pasitiki robotais ir žmonėmis. Tačiau, vertinant atskiras vaikų grupes pagal amžių, nustatyta, jog 3-ių metų amžiaus vaikai labiau pasitikėjo žmogumi, dėl emocinio prierašumo, o 7-ių – robotu, dėl suvokiamo patikimumo. Įvertinus pasitikėjimo lygmens pokytį žmogaus ir roboto klaidos atveju, nustatyta, jog vaikai vienodai greitai prarado pasitikėjimą abejais subjektais. Nors vaikai buvo linkę priskirti mažiau proto gebėjimų robotui nei žmogui, pasitikėjimo požiūriu skirtumai tarp abiejų subjektų nebuvo identifikuoti. Apibendrinant tyrimo rezultatus, galima teigti, jog suvokiamas veiksnumas teigiamai veikia vaikų pasitikėjimą robotu ir žmogumi.

PT taikoma ir rinkodaros, paslaugų, vartotojų elgsenos tyrimuose, kuriuose nagrinėjamas suvokiamas dirbtinių agentų veiksnumas. Li ir kt. (2023), remdamiesi PT teorija nustatė, jog suvokiamas protas glaudžiai sąveikauja su antropomorfizmu, vertinant DI ir žmogaus teikiamas telerinkodaros paslaugas. Tyrimas parodė, kad DI, kurie pasižymi didesniu antropomorfizmu, yra suvokiami, kaip galintys pasižymėti proto gebėjimais telerinkodaros paslaugose.

Remiantis PT, Söderlund'as (2024) siekė įvertinti suvokiamo proto vaidmenį vartotojų, paslaugų robotų ir paslaugą teikiančių žmonių sąveikoje bei nustatyti galimą poveikį vartotojų pasitenkinimui. Tyrimo rezultatai atskleidė, jog vartotojai yra linkę priskirti mažesnius proto gebėjimus paslaugų robotams nei žmonėms. Tačiau, šis reiškinys nėra vienareikšmiškas, o tam teigiamos įtakos turi polinkis antropomorfizuoti nežmogiškus subjektus. Be to, autorius pastebi, kad žmonės yra linkę dehumanizuoti paslaugą teikiančius žmones, šį reiškinį lemia daugelis veiksnių, pavyzdžiui, etiniai ar rasių skirtumai. Todėl, galima teigti, jog paslaugose, žmonės gali nepriskirti proto gebėjimų ir biologinėms būtybėms. Nustatyta, jog paslaugos palankumas ir orientacija į kliento poreikius turi įtakos vartotojo polinkiui priskirti didesnius proto gebėjimus paslaugų robotui. Taip pat, polinkis priskirti proto gebėjimus teigiamai veikia vartotojų pasitenkinimą ir skatina teigiamus atsiliepimus „iš lūpų į lūpas“ (angl. *word-of-mouth*, WOM). Galiausiai, suvokiamas protas turi įtakos bendram roboto vertinimui, kuomet šio reiškinio būvimas teigiamai veikia vartotojų požiūrį į šios technologijos taikymą paslaugose.

Zhao ir kt. (2025) tyrimas parodė, kad vartotojai, paslaugos nesėkmės atveju, yra linkę dažniau atleisti žmogiškomis savybėmis pasižymintiems paslaugų robotams, o tam įtakos turi suvokiamas roboto protas – patirtis ir veiksnumas. Rezultatai atskleidė, jog suvokiamas veiksnumas didina vartotojų pasitikėjimą ir užtikrina mažiau neigiamą reakciją į paslaugos klaidas.

Žinoma, jog polinkį priskirti moralinį veiksnumą paaiškina ir suvokiama moralė. Kita teorinė perspektyva, leidžianti geriau paaiškinti polinkį priskirti moralinį veiksnumą, yra suvokiamas

subjekto moralumas. Teigiama, jog moralė yra pagrindinis žmogaus tapatybės bruožas, kuris leidžia suvokti aplinką įvertinant ją „teisingumo“ ir „neteisingumo“ požiūriu (Luttrell et al., 2021). Moralė psichologijos mokslo srityje yra viena didžiausių mokslininkų susidomėjimą paskatinusių temų (Ramos et al., 2024). Bėgant metams, pastarųjų požiūris į moralės reiškinį kito, o vienas pirmųjų šiuolaikinės psichologijos mokslininkų, apibrėžusių moralę buvo Kohlberg'as (1969). Mokslininko teigimu, moralė susiformuoja dar ankstyvoje vaikystėje, kuomet ugdomas moralinių sprendimų suvokimas per (ne)sąžiningumo prizmę (Ramos et. al, 2024).

Evoliucionuojant požiūriui į moralę, psichologijos mokslo atstovai siekė atrasti sistemą, kuri padėtų suprasti kultūrinius, socialinius ir politinius skirtumus bei pateiktų vieningą atsakymą, kaip visuomenė suvokia šį reiškinį. 2004 m., mokslininkai Haidt'as ir Joseph'as pasiūlė moralinių pagrindų teoriją (angl. *moral foundations theory, MFT*), kuri tapo esmine moralės psichologijos dalimi, jungianti individualizuojančius ir saistančius moralės pagrindus. Individualizuojančiais moralės pagrindais siekiama moraliai apsaugoti ir rūpintis individualiais asmenimis, tuo tarpu saistančių pagrindų atveju fokusuojamasi į asmenų grupių, bendruomenių gerovę. Būtent šios dvi kategorijos – individualizuojančių ir saistančių moralės pagrindų – aktualios skirtingoms socialinėms ar kultūrinėms grupėms (Ramos et al., 2024).

Remiantis 2004 m. mokslininkų Haidt'o ir Joseph'o pasiūlyta moralinių pagrindų, moraliniai įvykiai ir juose dalyvaujantys subjektai gali būti vertinami atsižvelgiant į pagrindinius moralės principus, tokius, kaip rūpestis / žala, sąžiningumas / apgavystė, lojalumas / ištikimybė, autoritetas / subversija, tyrumas / nuopuolis (Banks, 2019). Rūpesčio ir sąžiningumo principai priskiriami individualizuojantiems moralės pagrindams, nes jie nusako individualaus asmens gerovę, o lojalumas, autoritetas ir tyrumas – saistantiems, pabrėždami grupės gerbūvį (Goenka & Thomas, 2024). Rūpesčio / žalos pagrindas nusako moralinius sprendimus, kurie pagrįsti suvokiamu gėrio ar žalos potencialu kito atžvilgiu (Telkamp & Anderson, 2022). Sąžiningumas / apgavystė apibrėžia suvokiamą moralinių sprendimų teisingumą. Lojalumas / ištikimybė nusako suvokiamą lojalumą savo grupei. Autoriteto / subversijos pagrindas apibrėžia socialinės tvarkos laikymosi pobūdį, pagarbos lygį tradicijoms. Tyrumas / nuopuolis dažniausiai siejamas su atsakinga elgsena, vengiančia psichologine ar fizine prasme netinkamų, degradacijai būdingų veiksmų (žr. **3 lentelė**).

3 lentelė. Moraliniai pagrindai

	Būdingos emocijos	Atitinkamos dorybės
Rūpestis/žala	Užuojauta aukai, pyktis ant nusikaltėlio	Rūpestingumas, gerumas
Sąžiningumas/apgavystė	Pyktis, dėkingumas, kaltė	Sąžiningumas, teisingumas, patikimumas
Lojalumas/ištikimybė	Pasididžiavimas grupe, pyktis ant išdavikų	Lojalumas, patriotizmas, pasiaukojimas
Autoritetas/subversija	Pagarba, baimė	Paklusnumas, pakantumas
Tyrumas/nuopuolis	Pasibjaurėjimas	Nuosaikumas, skaistumas, pamaldumas, švara

Šaltinis: Adaptuota pagal Haidt et al. (2013)

Su MPT teorija buvo pateiktas ir 30-ties punktų moralinių pagrindų klausimynas (angl. *Moral Foundations Questionnaire, MFQ*), kuris padeda įvertinti kiekvieną iš minėtų penkių moralės

pagrindų (žr. **3 priede**) (Graham et al., 2011). Šis klausimynas yra vienas plačiausiai naudojamų instrumentų moralinės psichologijos istorijoje (Zakharin & Bates), todėl jo patikimumas yra patvirtintas daugelyje mokslinių darbų. Klausimynas padalintas į dvi dalis, pirmoji jų nusako suvokiamų moralinių vertybių svarbą, tuo tarpu antroje dalyje siekiama suprasti, kaip vertinami konkretūs moraliniai sprendimai. Visi klausimai susieti su penkiais moraliniais pagrindais – rūpesčiu, sąžiningumu, lojalumu, autoritetu ir tyrumu.

Teigiama, jog žmonės skirtingai vertina ir suvokia moralines situacijas, o MPT nurodo, kad pastarieji yra linkę vertinti kitų moralinius sprendimus ir elgesį, remdamiesi šioje teorijoje pateiktais moraliniais pagrindais ir jų (ne)atitiktimi saviems moraliniams pagrindams (Telkamp & Anderson, 2022). Be to, visuomenėje egzistuoja skirtumai tarp individualių asmenų ir grupių bei polinkio priskirti pirmenybę MPT išskirtiems moraliniams pagrindams.

Šiuolaikinėje visuomenėje, kuomet moralinėse situacijose žmonių sąveiką papildo technologijos, šie skirtumai turi įtakos, kaip suvokiami ir vertinami ne tik žmonių, bet ir dirbtinių agentų veiksmai (Telkamp & Anderson, 2022). Tyrimai atskleidžia, jog žmonės yra linkę vertinti dirbtinių agentų moralę, remiantis savosiomis moralinėmis vertybėmis. Pavyzdžiui, Brailsford'as ir kt. (2024), remdamiesi MPT, atskleidė, jog sąveikaujant su DI moraliniu požiūriu svarbiose aplinkybėse, žmonės yra linkę remtis asmeninėmis moralinėmis vertybėmis, o didžiausias dėmesys skirtas rūpesčiui dėl numanomos žalos. Be to, šiam vertinimui įtakos turėjo ir suvokiamos DI techninės savybės. Mažesnės žmonių žinios apie DI galimybes, skatina žmonių polinkį antropomorfizuoti šią technologiją ir jai priskirti veiksnumą bei moralinę atsakomybę. Priešingu atveju, žmonės yra linkę suvokti DI kaip priemonę veiksams atlikti ir atsakomybę priskirti žmogui.

Suvokiamo moralumo svarba paslaugose grindžiama tuo, jog vartotojų ketinimai ir sprendimai reikšmingai lemiami turimų moralinių vertybių ir įsitikinimų (Ramos et al., 2024). Todėl, naujausiuose moksliniuose tyrimuose, MPT remiamasi ir vartotojų psichologijos tyrimuose. Pavyzdžiui, Ramos ir kt. (2024) tyrime, siekta išsiaiškinti kaip moraliniai pagrindai veikia vartotojų reakciją, emocijas, prosocialią ir tvarią elgseną. Tyrimas atskleidė, jog remiantis MPT galima geriau suprasti, kaip vartotojai reaguoja į produktus, paslaugas ar rinkodaros turinį. Goenka'o ir Thomas'o (2024) tyrimas praplėtė vartotojų psichologijos tyrimų lauką, atskleisdami, jog MPT gali padėti paaiškinti vartotojų moralinį naudingumą ir heterogeniškumą. Tačiau, vis dar trūksta tyrimų, kurie atskleistų, kaip suvokiamas moralumas veikia vartotojų atsaką į DI ir žmogaus moralinius sprendimus paslaugose.

Galima patvirtinti, jog PT yra svarbi suvokiamo dirbtinių agentų veiksnumo tyrimų dalis. Ši teorija plačiai taikoma ne tik psichologijos, žmonių elgsenos ir kitose susijusiose mokslo srityse, bet ir rinkodaros, paslaugų bei vartotojų elgsenos tyrimuose. Pastebima, kad didžioji dalis tyrimų analizuoja suvokiamą paslaugų robotų veiksnumą. Tačiau, mokslininkai pabrėžia, jog reikalinga toliau analizuoti suvokiamo DI veiksnumo poveikį vartotojų atsakui paslaugose, nes šios mokslo srities tyrimai dar tik įsibėgėja (Mariani et al., 2022). Be to, būtina tęsti ir plėsti tyrimus, analizuojant, kaip šis reiškinys skiriasi vartotojo sąveikoje su paslaugą teikiančiu žmogumi ir DI. Todėl, remiantis PT, šiame baigiamajame magistro projekte galima prognozuoti vartotojų polinkį priskirti proto savybes DI ir žmogui, kurios glaudžiai siejamos su suvokiamu veiksnumu. Kita vertus, aiškinant polinkį priskirti moralines savybes, galima teigti, jog MPT teorija paaiškina žmonių suvokiamą subjekto ar įvykio moralumą. Vartotojai, remdamiesi savo moraliniais pagrindais, yra linkę suvokti žmonių ir dirbtinių agentų moralinius sprendimus, kaip moraliai teisingus arba

neteisingus. Suvokiamas moralumas yra vienas iš minėtų, svarbiausių veiksnių, lemiančių polinkį priskirti moralinį veiksnumą (Banks, 2019). Apibendrinant, vartotojų polinkis priskirti moralinį veiksnumą priklausomai nuo subjekto tipo ir priimto sprendimo, bus aiškinamas remiantis PT ir MPT teorijomis. Teoriškai bus prognozuojama, kaip suvokiamas DI ir žmogaus protas bei moralė, veikia vartotojų polinkį priskirti moralinį veiksnumą DI ir žmogaus sprendimams.

2.3. Sprendimo palankumo ir polinkio priskirti moralinį veiksnumą ryšys

Kintančios paslaugų teikimo aplinkybės atskleidžia, jog vartotojai vis dažniau sąveikauja su DI grįstomis technologijomis. Šiandieninėje paslaugų aplinkoje, vartotojai susiduria ne tik su žmogaus, bet ir su DI priimtais sprendimais. Mokslinėje literatūroje identifikuoti trys vartotojų ir DI sąveikos paslaugose tipai, kurie nurodo, kaip ir koku mastu vartotojai sąveikauja su paslaugą teikiančiu DI ir žmogumi.

„DI palaikomas“ (angl. *AI-supported*) sąveikos tipas nurodo tiesioginę paslaugą teikiančių žmogiškųjų darbuotojų ir vartotojų sąveiką, kuomet DI pasitelkiamas tik kaip pagalbinė priemonė sprendimų priėmimui. „DI papildytas“ (angl. *AI-augmented*) sąveikos tipas nurodo tiesioginę DI ir vartotojų sąveiką, kuomet darbuotojai pasitelkia DI įrankius, siekiant papildyti tradicines paslaugų teikimo sąlygas, tačiau sprendimų priėmimas priklauso nuo žmogiškųjų darbuotojų. Galiausiai, „DI atliekamas“ (angl. *AI-performed*) sąveikos tipas nurodo atvejus, kai DI visapusiškai užima žmogiškojo darbuotojo vietą ir pats, savarankiškai priima sprendimus, pritaikydamas paslaugos patirtį realiuoju laiku (Xu et al., 2020).

Žinoma, jog vartotojai neretai susiduria su palankiais ir nepalankiais paslaugų sprendimais. Paslaugas atstovaujantys subjektai gali priimti vartotojų atžvilgiu teigiamus sprendimus, pavyzdžiui priimti ar patvirtinti tam tikras paslaugų teikimo sąlygas, ar procesus. Taip pat, egzistuoja situacijos, kuomet paslaugą teikiantys atstovai privalo reaguoti neigiamai – atmesti arba nepriimti vartotojui svarbių sąlygų ar procesų. Būtent šios aplinkybės apibūdina sprendimo palankumo sąvoką paslaugų kontekste (Yalcin et al., 2022).

Teigiama, jog sprendimo palankumas yra glaudžiai susijęs su sąžiningumu (Choi & Chao, 2024). Kitaip tariant, jei žmogus suvokia priimto sprendimo teisingumą, jis yra linkęs jį vertinti palankiai, kitu atveju, šis vertinimas yra neigiamas. Be to, suvokiamas sprendimo nesąžiningumas lemia tolimesnius žmonių ketinimus atrasti susijusias priežastis, kurios gali turėti įtakos sprendimo nepalankumui. Mokslininkų teigimu, žmonės, siekdami suprasti nepageidaujamo rezultato argumentus, yra linkę nagrinėti sprendimo priėmėjo savybes (Choi & Chao, 2024). Priimto sprendimo (ne)palankumas, gali turėti įtakos vartotojų reakcijai į paslaugos sandoryje dalyvaujančius subjektus (Yalcin et al., 2022).

Suvokiamam sprendimo (ne)palankumui bei reakcijai į jį, įtakos gali turėti paslaugą teikiantis subjektas. Šis reiškinyss gali būti grindžiamas remiantis Jones'o ir Davis'o, 1965 m. pasiūlyta atribucijos teorija (angl. *attribution theory*), kuomet vartotojai gali būti linkę paaiškinti priimto sprendimo priežastis, atsižvelgiant į vidinius ir išorinius veiksnus (Kelley, 1967). Yalcin ir kt. (2022) teigimu, sau palankius rezultatus vartotojai linkę priskirti vidinėms priežastims, t.y. remtis tuo, jog sprendimo palankumas yra lemiamas vidinių veiksnių, pavyzdžiui teigiamomis asmeninėmis savybėmis. Tuo tarpu nepalankius – išorinėms priežastims, vertinant priimto sprendimo priežastis, kurios gali būti lemiamos išorinių veiksnių, tokių, kaip sprendimo priėmėjas ar paslaugą teikianti įmonė (Yalcin et al., 2022).

Žinoma, jog sparčiai tobulėjant technologijoms, paslaugose sprendimus priima ne tik žmonės, bet ir DI. Todėl, naujausi tyrimai analizavo, kaip suvokiamas sprendimo (ne)palankumas priklauso nuo paslaugą teikiančio subjekto tipo – DI ir žmogaus (Choi & Chao, 2024; Yalcin et al., 2022). Socialinės psichologijos mokslo atstovų Choi ir Chao (2024) tyrimu siekta įvertinti žmonių reakciją į DI ir žmogaus nepalankius sprendimus bei nustatyti pastarųjų reakciją lemiančius veiksnius. Rezultatai atskleidė, jog žmonės yra linkę mažiau neigiamai reaguoti į DI nepalankius sprendimus, o tam įtakos turi suvokiamas DI teisingumas ir nešališkumas. Žmonių nuomone, žmonės, palyginti su DI, šiomis savybėmis gali pasižymėti dažniau. Tačiau, priminus, jog DI gali pateikti neobjektyvius rezultatus, dėl taikomo išvesties šališkumo (Akter et al., 2022), DI ir žmogaus sprendimai teisingumo požiūriu buvo vertinami identiška. Todėl, tyrimas patvirtina, jog žmonės gali skirtingai suvokti DI ir žmogaus priimtus (ne)palankius sprendimus ir yra linkę į juos reaguoti nevienodai.

Tuo tarpu rinkodaros mokslo srityje šią temą nagrinėjo mokslininkai Yalcin ir kt. (2022), siekdami suprasti, kaip suvokiamas sprendimo (ne)palankumas veikia vartotojų reakciją (teigiamai ar neigiamai) į žmogaus arba algoritmo pasiūlytą sprendimą. Tyrimas parodė, kad vartotojai geriau reaguoja į žmogaus pateiktą vartotojui palankų sprendimą, o tai lemiamą vidinę atribuciją, siekiant padidinti savo vertę. Nepalankaus sprendimo atveju, abu subjektai vertinami vienodai prastai, o tai grindžiama polinkiu paaiškinti sprendimo priežastis, remiantis išoriniais veiksniais, kaip paslaugos teikėjo nedėmesingumas, šališkumas, subjektyvumas (Yalcin et al., 2022).

Esamų tyrimų apžvalga atskleidžia, jog sprendimo palankumo tema, vertinant skirtingus sprendimo priėmėjus, ypač dirbtinius agentus, yra pakankamai nauja (Yalcin et al., 2022). Tačiau, mokslinių darbų išvados rodo, jog žmonės yra linkę skirtingai reaguoti į žmogaus ir dirbtinių agentų pateiktus (ne)palankius sprendimus. Be to, teigiama, jog (ne)palankus požiūris į DI ir žmogaus sprendimus gali lemti vartotojų elgesio ketinimus (Jain et al., 2024), todėl ši tema aktuali rinkodaros mokslo atstovams.

Apibendrinus, šiame projekte sprendimo (ne)palankumas bus siejamas su vartotojo lūkesčių (ne)atitikimu. Apžvelgus esamus empirinius tyrimus, galima daryti prielaidą, jog sprendimo (ne)palankumas bei sprendimo priėmėjo tipas gali būti aktualus ir moraliniu požiūriu svarbiose paslaugų teikimo aplinkybėse. Atsižvelgiant į mokslininkų išvadas, galima manyti, jog suvokiamas sprendimo (ne)palankumas gali veikti vartotojų polinkį priskirti veiksnio ir moralumo savybes paslaugą teikiančiam subjektui. Ši prielaida grindžiama mokslinių tyrimų rezultatais, kurie atskleidžia, jog vartotojų reakcija gali priklausyti nuo suvokiamų dirbtinių agentų ir žmogaus savybių, pavyzdžiui šališkumo (Choi & Chao, 2024).

Esama literatūra atskleidžia egzistuojančius vartotojų reakcijos skirtumus, vertinant žmonių ir dirbtinių agentų sprendimų (ne)palankumo įtaką. Vis dėlto, nėra žinoma, kokią įtaką suvokiamas sprendimo (ne)palankumas gali turėti vartotojų reakcijai į sprendimo priėmėją (DI ir žmogų) moraliai abejotinoje paslaugų teikimo aplinkybėse. Be to, nėra aišku, kaip sprendimo (ne)palankumas gali veikti vartotojų polinkį priskirti moralinį veiksnumą vienam iš paslaugą teikiančių subjektų.

Todėl, atsižvelgiant į mokslinės literatūros spragas ir egzistuojančias išvadas, šiame baigiamajame magistro projekte bus remiamasi atribucijos teorija, siekiant suprasti, kaip suvokiamas priimto moralinio sprendimo (ne)palankumas veikia vartotojų polinkį priskirti moralinį veiksnumą identiškiems DI ir žmogaus sprendimams paslaugose.

2.4. Suvokiamo sąžiningumo konceptualizacija

Mokslininkai sutaria, jog vis dar nėra vienijančios sąžiningumo reiškinio apibrėžties (Shulner-Tal et al., 2023). Anot Wang ir kt. (2022), „sąžiningumas yra santykinė socialinė sąvoka, todėl absoliučia prasme sąžiningumo nėra.“ (p. 9). Žvelgiant plačiaja prasme, sąžiningumas siejamas su subjekto priimtu racionalių sprendimų veikti sąžiningai, remiantis visuomenėje ar konkrečioje grupėje paplitusiomis teisingumo normomis (Helberger et al., 2020). Sąžiningumas siejamas su nešališku ir vienuodu gėrybių bei teisių paskirstymu, remiantis asmenų veiklos rezultatais bei poreikiais (Ochmann et al., 2024).

Teigiama, jog suvokiamas sąžiningumas yra subjektyvus procesas. Žmonės yra linkę suvokti sąžiningumą remiantis vidine intuicija ir afektine informacija (Narayanan et al., 2024). Suvokiamas sąžiningumas priklauso nuo daugybės kintamųjų, kaip kultūra, socialinė aplinka, skirtingi kontekstai bei individualios asmens savybės (Helberger et al., 2020). Be to, šis procesas gali būti susijęs ir su tam tikroje padėtyje dalyvaujančiais subjektais bei jų charakteristikomis. Nors suvokiamas sąžiningumas yra lemiamas daugelio susijusių veiksnių ir gali būti interpretuojamas priklausomai nuo kontekstinių aplinkybių, neretai ši sąvoka siejama su moraliai teisingų principų laikymusi bei šių principų atitiktimi socialinėms vertybėms (Shulner-Tal et al., 2023).

Paslaugų kontekste, suvokiamas sąžiningumas apibrėžiamas, kaip vartotojo suvokiamas paslaugų sandoryje dalyvaujančio subjekto elgsenos teisingumo laipsnis (Wu et al., 2024). Vartotojas savo sąmonėje vertina priimto sprendimo ir rezultato pagrįstumą, tinkamumą, ir teisingumą, o tai yra esminis etikos, ir moralės principas, kuris formuoja vartotojo, ir įmonės tarpusavio santykius (Mathur & Sarin, 2020). Be to, akcentuojama, jog suvokiamas nesąžiningumas turi įtakos vartotojo reakcijai į kitus žmones, grupes ar įmones (Mathur & Sarin, 2020) bei gali sukelti kitas neigiamas pasekmes, pavyzdžiui pakenkti ilgalaikiai įmonių reputacijai (Weith & Matt, 2022). Panašiai teigia Wu ir kt. (2024), pabrėždami, jog vartotojų suvokiamas nesąžiningumas gali itin neigiamai veikti vartotojų atsaką trumpuoju (pavyzdžiui, pasitenkinimą) ir ilguoju (pavyzdžiui, lojalumą) periodais.

Wu ir kt. (2024) teigimu, 2001 m., Folger'io ir Cropanzano pasiūlyta sąžiningumo teorija (ST) (angl. *fairness theory, FT*) yra esminė sąžiningumo tyrimų dalis, plačiai paplitusi įvairiose mokslo srityse. Ši teorija nurodo, jog sąžiningumas siejamas su suvokiama asmens atskaitomybe priimto sprendimo atžvilgiu (Liu & Mendez, 2022). Suvokiamos atskaitomybės vertinimai apima tris kontrafaktus, kurie apibūdina minčių modeliavimo procesą, kvestionuojantį praeities įvykių alternatyvas. Kitaip tariant, kontrafaktai nurodo sąlyginius teiginius su klaidingais antecedentais, siekiant įvertinti „tai, kas galėjo būti“, susiklosčius konkrečioms aplinkybėms (De Brigard, 2023). Akcentuojama, jog vertinant subjekto atskaitomybę remiamasi trimis hipotetiniais klausimais, kuriais siekiama įvertinti priimto sprendimo (ne)teisingumą (Liu & Mendez, 2022). Suvokiamo neteisingumo sąlyga apima teigiamus atsakymus į žemiau nurodytus kontrafaktus:

1. Kontrafaktas „norėjo“ (angl. *would*). Šiuo teiginiu siekiama įvertinti patirtos žalos mastą, kvestionuojant „ar būtų buvę geriau, jei rezultatas, ar procedūra būtų buvusi kitokia?“;
2. Kontrafaktas „galėjo“ (angl. *could*). Šiuo teiginiu kvestionuojama elgsena ir pasirinkti veiksmai, svarstant „ar galėjo būti pasielgta kitaip?“;
3. Kontrafaktas „turėjo“ (angl. *should*). Galiausiai, siekiama įvertinti atitikmenį moralės ir etikos normoms, keliant klausimą „ar turėjo būti pasielgta kitaip?“ (Liu & Mendez, 2022, cit. iš Colquitt, et al., 2005).

Todėl, remiantis ST, suvokiamas sąžiningumas apibūdinamas trejomis sisteminėmis sritimis, kurios sieja bendro gerbūvio, individualaus elgesio ir moralės bei etikos normų laikymosi vertinimą. Be to, svarbu paminėti, jog kontrafaktai „norėjo“, „turėjo“ ir „galėjo“ neturi nuoseklios tvarkos ir gali būti išdėstyti atsitiktinai. Akcentuojama, jog egzistuoja atvejai, kuomet atskaitomybės vertinimas įvyksta anksčiau, nei tampa žinoma apie galutinį poveikį (Liu & Mendez, 2022). Atskaitomybės padidėjimas yra siejamas su suvokiama neteisybe ir priskiriamu neteisingumo laipsniu. Ši sąlyga paslaugų kontekste išryškėja, kuomet vartotojų patirtis neatitinka turimų sąžiningumo standartų, o tai lemia vartotojų suvokiamą paslaugų subjekto neteisingą elgseną (Wu et al., 2024). Be to, atskaitomybės padidėjimas susijęs su polinkiu priskirti kaltę sąveikoje dalyvaujančiam subjektui (Jung et al., 2022). Remiantis sąžiningumo teorija, būtina suvokiamo sąžiningumo sąlyga nurodo tai, jog įvykus neigiamam rezultatui, turi būti subjektas, kuris bus laikomas atlikusiu tokį rezultatą lėmusį veiksmą (Jung et al., 2022). Paslaugose vartotojai yra linkę vertinti savų išteklių sąnaudas ir palyginti jas su paslaugų sandoryje dalyvaujančiais subjektais, siekiant įvertinti, kas gali būti kaltas, dėl paslaugos metu įvykusio neteisingo rezultato (Scholl-Grissemann et al., 2022). Apibendrinant galima teigti, jog remiantis ST, vartotojų suvokiamam nesąžiningumui įtakos turi suvokiama žala, neproporcingai mažesnis paslaugų teikėjo indėlis bei elgsenos neatitikimas etikos ir moralės normoms, o tai turi tiesioginės įtakos vartotojų vertinimams, ir tolimesnei reakcijai (Scholl-Grissemann et al., 2022).

Sparčiai augantis DI technologijų diegimas su paslaugomis susijusių sprendimų procese, kelia vis didesnę akademinės bendruomenės susidomėjimą, kvestionuojant, kokį vaidmenį vartotojų ir DI sąveikoje atlieka suvokiamas DI sąžiningumas (Narayanan et al., 2024). Naujausioje mokslinėje literatūroje skiriamas ypatingas dėmesys suvokiamo DI sąžiningumo konceptualizavimui, nes vis daugiau svarbių sprendimų įvairiuose kontekstuose yra priimami DI algoritmų (Shulner-Tal et al., 2023).

Nors egzistuoja skirtingi požiūriai į DI sąžiningumą, bendrąja prasme ši sąvoka apima priimamų sprendimų teisingumą, nekeliant diskriminacinių ar nevienodų pasekmių individualiems asmenims ir jų grupėms (Starke et al., 2022). Egzistuoja kelios suvokiamo DI sąžiningumo tyrimų kryptys, vienoje jų fokusuojamasi į matematiškai apibrėžiamą DI sąžiningumą, kitoje gilinamasi į filosofinę ir socialinę žmogaus suvokiamo sąžiningumo koncepciją (Starke et al., 2022). Pastarojoje mokslinių tyrimų grupėje, didžioji dalis suvokiamo DI sąžiningumo tyrimų buvo atlikti žmogaus ir kompiuterio sąveikos mokslo srityje, siekiant suprasti, kaip suvokiamas DI sąžiningumas sistemų projektuotojų, sprendimų priėmėjų ir kitų suinteresuotų šalių požiūriu (Narayanan et al., 2024). Tačiau, autoriai pastebi mokslinės literatūros spragas, analizuojant, kaip suvokiamas DI sąžiningumas, kai technologijos priimti sprendimai tiesiogiai veikia galutinius vartotojus. Panašiai teigia Weith ir Matt'as (2022), pabrėždami, jog vis dar menkai nagrinėta vartotojo perspektyva DI sąžiningumo suvokimo tyrimuose. Iki šiol atlikti tyrimai atskleidžia, jog suvokiamas sąžiningumas yra viena iš svarbiausių vartotojų požiūrio ir patirties sąveikaujant su DI dedamųjų (Choung et al., 2023), todėl tampa svarbu toliau nagrinėti, kaip šis reiškinys suvokiamas vartotojų požiūriu.

Siekiant paaiškinti suvokiamo DI sąžiningumo reiškinį, iki šiol atlikti tyrimai dažnai rėmėsi 2001 m. mokslininko Colquitt'o ir kt. pasiūlytomis keturiomis sąžiningumo dimensijomis. Remiantis šiuo požiūriu, sąžiningumo vertinimas priklauso nuo procedūrinio, paskirstomojo, tarpasmeninio ir informacinio teisingumo suvokimo (Ochmann et al., 2024). Remiantis Weith'u ir Matt'u (2022), kiekvienai dimensijai priskiriamos taisyklės, kuriomis remiantis vertinamas DI sąžiningumas paslaugų kontekste. Žemiau pateikiami paskirstomojo, procedūrinio, tarpasmeninio ir informacinio teisingumo dimensijas detalizuojantys aprašymai (žr. 4 lentelė).

4 lentelė. Suvokiamo DI sąžiningumo dimensijos ir taisyklės

Sąžiningumo dimensija	DI sąžiningumo taisyklė	DI sąžiningumo taisyklės aprašymas
Procedūrinė	Proceso kontrolė	Procesas leidžia daryti įtaką paslaugų teikimo etapams
	Sprendimų kontrolė	Procesas suteikia galimybę priimti arba atmesti rezultatus
	Koreguojamumas	Procesas leidžia atšaukti ir keisti paslaugos etapus
	Pritaikytas nuoseklumas	Procesas užtikrina nuoseklumą tarp klientų, tačiau laikui bėgant, DI užtikrina prisitaikymą prie vieno kliento
	Numatymas	Procesas leidžia numatyti būsimą paklausą, kurią lemia dirbtinis intelektas
	Šališkumo slopinimas	Procesas slopina šališkumą DI ir kitų pelno siekiančių šalių labui
	Tikslumas	Procesas užtikrina susijusios informacijos tikslumą
	Etiškumas	Procesas užtikrina, kad būtų laikomasi asmens moralės ir vertybių
	Kompetencija	Procesas demonstruoja DI kompetenciją
	Saugumas	Procesas užtikrina duomenų saugumą ir apsaugą
Paskirstomoji	Duomenimis pagrįstas teisingumas	Rezultatas paskirstomas remiantis pateiktų duomenų kiekiu ir iš jų gautomis įžvalgomis
	Individualumas	Rezultatas personalizuojamas kiekvienam klientui
	Lygybė	Rezultatas yra vienodai lygus visiems klientams
	Alternatyvos	Rezultatas apima klientams aktualias alternatyvas
	Nepriklausomybė	Rezultatas nepriklauso nuo trečiųjų šalių interesų
Tarpasmeninė	Pagarba	Rezultatas yra dėmesingas ir mandagus
	Tinkamumas	Rezultatas yra tinkamas ir pakankamas
Informacinė	Tikrumas	Informacija apie procesą ir rezultatą yra sąžininga ir skaidri
	Pagrindimas	Proceso ir rezultato paaiškinimai yra išsamūs, ir tikėtini

Šaltinis: Adaptuota pagal Weith & Matt, 2022

Procedūrinis teisingumas apima procedūrų vertinimą, atsižvelgiant į išteklių paskirstymą ir sprendimo rezultata, t.y. vertinamos sprendimų priėmimo procedūros (Ochmann et al., 2024). Procedūrinis teisingumas apibrėžiamas remiantis aštuoniomis teisingumo taisyklėmis, įskaitant proceso kontrolę, sprendimų kontrolę, korektiškumą, nuoseklumą, šališkumo slopinimą, tikslumą, atstovavimą ir etiškumą (Weith & Matt, 2022). Šiuo atveju teisingumas suvokiamas, kai procedūros laikomos skaidrios, nešališkos, tikslingos ir išsprendžiamos, kitaip tariant, procedūros gali būti keičiamos, dėl išreikšto žmonių poreikio (Narayanan et al., 2024).

Paskirstomasis teisingumas nurodo gautų sprendimo rezultatų vertinimą. Priimtų sprendimų rezultatų vertinimas grindžiamas išteklių paskirstymo taisyklėmis, kurios buvo taikytos priimtų sprendimų procese (Ochmann et al., 2024). Šiuo atveju, suvokiamam teisingumui įtakos turi priimtų sprendimų atitikmuo išteklių paskirstymo taisyklėms, t.y. teisingumui ir lygybei (Narayanan et al., 2024). Teisingumo normos siejamos su išteklių paskirstymo procesu, kuomet atsižvelgiama į asmeninį indėlį, numanomą nuopelną, tuo tarpu lygybės normos apibūdina vienodą išteklių paskirstymą visiems asmenims ir jų grupėms (Narayanan et al., 2024).

Tarpasmeninis teisingumas siejamas su tarpasmeninės sąveikos teisingumu, vertinant elgsenos su žmonėmis mandagumą, orumą ir pagarbą (Ochmann et al., 2024). Be to, ši dimensija siejama ir su procedūriniu teisingumu, nes tarpasmeniniuose santykiuose vertinama reakcija į sprendimo priėmimo procesus ir refleksijos pobūdis į sprendimų rezultatus (Narayanan et al., 2024).

Galiausiai, **informacinis teisingumas** vertinamas atsižvelgiant į priimtų procedūrų paaiškinimą ir gautų rezultatų pagrindimą (Ochmann et al., 2024). Suvokiamam informaciniam teisingumui įtakos turi informacijos savalaikiškumas ir pagrįstumas (Ochmann et al., 2024). Šiuo atveju informacijos teisingumas siejamas su teisingumo ir pagrindimo taisyklėmis, pagal kurias reikalaujama sąžiningų ir išsamių paaiškinimų (Weith & Matt, 2022).

Anot Narayanan'o ir kt. (2024), Colquitt'o pasiūlytos sąžiningumo dimensijos yra itin aktualios suvokiamo DI sąžiningumo požiūriu, nes DI charakteristikos ir su tuo kylantys iššūkiai yra susiję su kiekviena dimensija. Todėl, vartotojų suvokiamas DI sąžiningumas gali būti vertinamas nagrinėjant atskiras sąžiningumo dimensijas arba jų tarpusavio sąsajas, priklausomai nuo paslaugų konteksto, technologijos tipo, jos taikymo aplinkybių bei sąveikos pobūdžio. Šiame projekte, suvokiamas sąžiningumas bus analizuojamas remiantis dvejomis sąžiningumo dimensijomis: procedūrine ir tarpasmenine. Procedūrine dimensija bus siekiama įvertinti, kaip suvokiamo sąžiningumo reiškinys siejasi su tyrimui aktualiu suvokiamu DI ir žmogaus veiksmu. Tuo tarpu tarpasmeninė dimensija padės įvertinti suvokiamo sąžiningumo ir suvokiamo moralumo santykį. Bendrai, procedūrinė ir tarpasmeninė dimensijos padės įvertinti vartotojų suvokiamo sąžiningumo ir polinkio priskirti moralinį veiksmumą DI ir žmogaus sprendimams sąsajas.

Autorių teigimu, suvokiamam DI sąžiningumui įtaką gali daryti daugelis susijusių veiksnių, pavyzdžiui **sociodemografiniai rodikliai**. Van Berkel'is ir kt. (2021) nustatė, kad suvokiamam technologijos sąžiningumui įtakos turi žmogaus lytis. Tyrimas atskleidė, jog moterys yra linkę priskirti mažesnę sąžiningumo lygį ir mažiau pasitikėti DI priimtais sprendimais. Be to, įtakos turi amžius ir išsilavinimo lygmuo, kuomet jaunesni žmonės yra jautresni suvokiamo sąžiningumo požiūriu (Grgić-Hlača et al., 2022), o aukštesnis išsilavinimas lemia griežtesnius sąžiningumo vertinimus bei mažesnę suvokiamą sąžiningumą (van Berkel et al., 2021). Tuo tarpu technologinis

raštingumas vertinamas prieštaringai, vieni autoriai teigia, jog aukštas raštingumas lemia žemesnį suvokiamą sąžiningumą (Schoeffer et al., 2021a), kiti – didesnį (Araujo et al., 2020).

Suvokiamas sąžiningumas siejamas su informuotumu apie DI skaidrumą, pateikiant **priimtų sprendimų paaiškinimus**. Angerschmid ir kt. (2022) nustatė, kad DI sprendimų paaiškinimai reikšmingai didina suvokiamą DI sąžiningumą. Autoriai nepastebėjo skirtingų paaiškinimų tipų (pavyzdžiais pagrįsto paaiškinimo ir funkcijų svarba pagrįsto paaiškinimo) poveikio šiam reiškiniui. Įrodyta, jog DI priimtų sprendimų paaiškinimas teigiamai veikia vartotojų atsaką, pavyzdžiui pasitikėjimą (Shin, 2021). Tačiau, įtakos turi ir pateikiamų paaiškinimų pobūdis, kuomet jų beprasmisškumas gali sukelti neigiamą reakciją (Angerschmid et al., 2022). Nors DI skaidrumas teigiamai veikia suvokiamą sąžiningumą, Ebrahimi ir kt. (2024) pažymi, kad šios sąlygos gali būti veikiamos žmonėms būdingo riboto racionalumo ir kognityvinių šališkumų. Todėl, galima teigti, jog net ir esant DI skaidrumui, vartotojų žmogiškosios savybės gali paveikti šį suvokiamo sąžiningumo veiksnį.

Mokslininkai teigia, jog suvokiamam DI sąžiningumui įtakos gali turėti ir **priimto sprendimo palankumas**. Choung ir kt. (2023) atskleidė, jog sprendimo palankumas skatina žmonių polinkį priskirti teigiamus atributus sprendimo priėmėjui ir tokiu atveju DI suvokiamas, kaip labiau sąžiningas. Panašiai nurodo Banks ir kt. (2022), teigdami, jog padidėjusį suvokiamą DI sąžiningumą lemia priimtas žmogui palankus, teigiamas sprendimas. Šią įžvalgą patvirtina Choi ir Chao (2024), jog suvokiamas DI sąžiningumas didėja priklausomai nuo priimto sprendimo rezultato palankumo. Tačiau, Yalcin ir kt. (2022) nurodo priešingai, kad rinkodaroje, vartotojų suvokiamas (ne)sąžiningumas nėra priklausomas nuo sprendimo (ne)palankumo.

Be to, nustatyta, jog suvokiamas sąžiningumas gali priklausyti nuo **skirtingų sprendimo priėmėjų tipų – DI arba žmogaus**. Šis reiškinys gali būti grindžiamas tuo, jog požiūris į informacijos šaltinį lemia informacijos patikimumo ir kokybės vertinimus (Lee, 2018). Žmonių suvokiamas DI ir žmogaus sąžiningumas plačiai nagrinėtas įvairiose visuomeninėse (Helberg et al., 2020; Lee & Rich, 2021, Lee, 2018) bei organizacinėse srityse (Qin et al., 2023; Noble et al., 2021), o pastaruoju metu atlikti moksliniai darbai atskleidžia, jog ši tema pradėta analizuoti rinkodaros ir paslaugų tyrimuose (Wu et al., 2024). Vis dėlto, naujausi empiriniai duomenys atskleidžia, jog šiuo požiūriu egzistuoja prieštaringos mokslinių tyrimų išvados. Dalis autorių teigia, jog žmonės yra linkę manyti, kad DI pasižymi didesniu sąžiningumu nei žmogus. Šią išvadą patvirtina Helberger ir kt. (2020) tyrimo rezultatai, rodantys, jog DI suvokiamas, kaip sąžiningesnis sprendimų priėmėjas, o kai kurie žmonės yra linkę idealizuoti DI galimybes sprendimų priėmimo procese. Qin'as ir kt. (2023) atskleidė, kad darbuotojų veiklos vertinimo požiūriu DI suvokiamas, kaip galintis priimti sąžiningesnius sprendimus nei žmogiškasis vadovas. Wu ir kt. (2024) tyrimas parodė, kad vartotojai yra linkę vertinti DI paslaugų robotus, kaip labiau sąžiningus, lyginant su paslaugą teikiančiais žmonėmis. Šios išvados ypatingai išryškėjo paskirstomojo ir procedūrinio sąžiningumo požiūriu, kuomet vartotojai buvo linkę labiau pasitikėti ir tikėti DI priimtų sprendimų, ir rezultatų teisingumu tokiose paslaugų srityse, kaip restoranai, viešbučiai, ir medicina. Kiti autoriai teigia priešingai, jog žmonės suvokiami, kaip pasižymintys didesniu sąžiningumu nei DI. Lee (2018) tyrimas parodė, kad DI sprendimai suvokiami, kaip mažiau sąžiningi nei žmogaus. Darbuotojų atrankos kontekste, žmonėms priskiriamas didesnis sąžiningumas nei DI, ypač procedūrinio ir tarpasmeninio sąžiningumo požiūriu (Noble et al., 2021). Todėl, apibendrinus, suvokiamas sąžiningumas DI ir žmogaus atžvilgiu vertinamas nevienodai. Šis reiškinys gali būti priklausomas nuo skirtingų veiksnių, kaip suvokiamas subjekto šališkumas, pasitikėjimo lygmuo, priimto sprendimo sudėtingumas ir kompleksisškumas.

Suvokiamo sąžiningumo reiškiniui įtakos turi suvokiamas **DI ir žmogaus šališkumas**. Žmonės yra linkę manyti, jog DI sprendimų priėmimo procese pasižymi mažesniu tarpasmeniniu šališkumu nei žmogus, o tai turi įtakos suvokiamam didesniai DI sąžiningumui (Qin et al., 2023). Panašias išvadas pateikia Bankins ir kt. (2022) tyrimo rezultatai, akcentuojant, jog žmonės suvokia DI, kaip mažiau šališkus ir labiau objektyvius, todėl yra linkę manyti, jog DI pasižymi didesniu sąžiningumu nei žmogus. Panašiai teigia Schoeffer'is ir kt. (2021b), jog žmonių suvokiamas DI nešališkumas didina suvokiamą DI priimtų sprendimų sąžiningumą, lyginant su žmogumi. Šios išvados grindžiamos suvokiamu DI objektyvumu, žmonių subjektyvumu ir emocionalumu. Choi ir Chao (2024) atskleidė, kad suvokiamas DI ir žmogaus šališkumas neigiamai veikia suvokiamą sąžiningumą. Informuojant apie galimą DI šališkumą, autorių teigimu, „galima susilpninti DI priskiriamą teisingumo aureolę, palyginti su žmonėmis.“ (Choi & Chao, 2024, p. 15).

Tyrimai atskleidžia, kad suvokiamas DI ir žmogaus sąžiningumas gali priklausyti nuo **pasitikėjimo lygmens**. Pavyzdžiui, Lee ir Rich (2021) atskleidė, kad pasitikėjimas žmonių priimamais sprendimais medicinos srityje veikia polinkį priskirti didesnį sąžiningumą žmonėms, bet ne DI. Tuo tarpu žmonės, kurie nepasitiki žmogaus sprendimais yra linkę manyti, jog DI gali priimti sąžiningesnius sprendimus.

Suvokiamas DI ir žmogaus sąžiningumas yra siejamas su **priimto sprendimo sudėtingumu ir kompleksiskumu**. Araujo ir kt. (2020) nustatė, jog suvokiamas sąžiningumas siejamas su mažo / didelio poveikio sprendimais, analizuojant tokias sritis, kaip teisingumas, sveikata ir žiniasklaida. Didelio poveikio sprendimuose (ypač teisingumo ir sveikatos srityse) DI buvo priskiriamas didesnis sąžiningumo lygmuo, tuo tarpu mažo poveikio sprendimuose suvokiamas sąžiningumas tarp DI ir žmogaus nesiskyrė. Tačiau, kuomet sprendimas reikalauja išskirtinai žmogiškų savybių (pvz. emocinių) bei įgūdžių, žmonės yra linkę manyti, kad DI pasižymi mažesniu sąžiningumu nei žmogus (Lee, 2018). Be to, itin kompleksinio požiūrio reikalaujančiuose sprendimuose, kurie apima daugybę veiksmų ir etapų, didesnis sąžiningumas priskiriamas žmogui, lyginant su DI (Gupta et al., 2022). Todėl, galima patvirtinti, jog suvokiamas DI ir žmogaus sąžiningumas gali būti lemiamas užduoties sudėtingumo, ir kompleksiskumo (Narayanan et al., 2024). Šis veiksnys itin aktualus moralinėse situacijose, nes moraliniu požiūriu svarbūs sprendimai dažnai būna sudėtingi (Dietvorst & Bartels, 2022) ir turi didelę įtaką vartotojų reakcijai bei atsakui.

Mokslinė literatūra atskleidžia, jog žmonės skirtingai vertina DI ir žmogaus sąžiningumą. Šį reiškinį veikiamas individualių asmens savybių, priimto sprendimo (ne)aiškumo, (ne)palankumo, (ne)sudėtingumo ir (ne)kompleksiskumo, taip pat įtakos turi (ne)pasitikėjimas sprendimo priėmėju. Nustatyta, jog rinkodaros mokslo srityje ši tema dar tik pradėta nagrinėti, todėl trūksta žinių apie tai, kaip šie veiksniai lemia suvokiamą DI ir žmogaus sąžiningumą paslaugose. Be to, nėra žinoma, kaip suvokiamas DI ir žmogaus sąžiningumas moraliniu požiūriu svarbiose paslaugų aplinkybėse.

Apibendrinant, šiame projekte suvokiamas DI ir žmogaus sprendimų sąžiningumas apibrėžiamas remiantis sąžiningumo teorija (ST), kaip dvidimensinis reiškinys, apimantis procedūrinį ir tarpasmeninį teisingumą. Įvertinus minėtas sąžiningumo dimensijas bus siekiama suprasti, kaip vartotojų polinkis priskirti moralinį veiksnumą veikia suvokiamą DI arba žmogaus sąžiningumą, moraliai abejotinosiose paslaugų teikimo aplinkybėse. Todėl, šis projektas praplės esamą supratimą apie suvokiamą DI ir žmogaus sąžiningumo reiškinį paslaugų kontekste.

2.5. Suvokiamo sąžiningumo ir pasitenkinimo paslaugomis ryšys

Apžvelgus suvokiamo sąžiningumo reiškinį, tampa svarbu suprasti, kokią įtaką suvokiamas sąžiningumas gal turėti vartotojų pasitenkinimui paslaugomis, nes būtent vartotojų atsakas yra esminis šiuolaikinės rinkodaros siekis (Brandtner et al., 2021). Teigiama, jog vartotojų pasitenkinimas yra būtina ilgalaikių santykių su vartotojais kūrimo ir stiprinimo sąlyga, nes ši vartotojų atsako forma yra esminė elgsenos lojalumo prielaida (Leclercq-Machado et al., 2022).

Vartotojų pasitenkinimas apibrėžiamas, kaip reakcijos į tam tikrą tikslą (pvz.: vartojimo patirtį, lūkesčius), kuri pasireiškia konkrečiu laiko periodu (pvz.: po produkto/paslaugos pasirinkimo), vertinimo proceso rezultatas (Brandtner et al., 2021). Paslaugų kontekste, vartotojai vertina paslaugos teikimo rezultatus, lygindami juos su suvokiamais paslaugos rezultatais (Jimenez Mori, 2021). Kitaip tariant, vartotojų pasitenkinimas yra veikiamas kognityvinių ir emocinių būsenų, todėl paslaugos rezultatai vertinami remiantis išankstiniais lūkesčiais bei afektiniu požiūriu į paslaugą, ir aplinkos dirgiklius (Ruan & Mezei, 2022). Teigiama, kad (ne)patvirtinti lūkesčiai daro įtaką vartotojų (ne)pasitenkinimui paslaugomis (Wang et al., 2021).

Žinoma, jog lūkesčiai gali priklausyti nuo ankstesnės vartojimo patirties bei bendro suvokimo, ar rezultatai yra sąžiningi (Nicolitsas, 2024). Atsižvelgiant į tai, jog paslaugos pasižymi neapčiuopiamumu, vartotojai vertina paslaugos rezultatus pagal tai, kaip sąžiningai su jais elgiamasi (Kim et al., 2023). Vartotojai reaguoja į paslaugą teikiančio subjekto indėlį sandorio metu, siekdami įvertinti situacijos teisingumą (Kim et al., 2023). Moksliniai duomenys atskleidžia, jog suvokiamas (ne)sąžiningumas gali sukurti situacijos (ne)teisingumo įspūdį ir daryti įtaką vartotojų (ne)pasitenkinimui paslaugomis (Wang et al., 2021).

Esamų tyrimų apžvalga atskleidžia, jog suvokiamo sąžiningumo ir pasitenkinimo paslaugomis tema plačiausiai nagrinėta analizuojant suvokiamo kainų sąžiningumo reiškinį įvairiuose paslaugų sektoriuose. Pavyzdžiui, Alzoubi'o ir kt. (2020) tyrimas parodė, kad suvokiamas kainų sąžiningumas telekomunikacijų paslaugose teigiamai veikia vartotojų pasitenkinimą. Be to, nustatyta, kad suvokiamas įmonės sąžiningumas po paslaugų sutrikimo turi įtakos vartotojų pasitikėjimui, kuris siejamas su ilgalaikiu pasitenkinimu. Autoriai pabrėžia, kad vartotojų suvokiamas sąžiningumas yra glaudžiai susijęs su pasitenkinimu paslaugomis, o tai prasmingai prisideda prie augančio vartotojų lojalumo ir ilgalaikės verslo sėkmės.

Setiawan'as ir kt. (2020) analizavo vartotojų suvokiamą Indonezijos oro linijų bendrovių kainų sąžiningumą ir šio reiškinio poveikį pasitenkinimui. Tyrimas atskleidė, kad suvokiamas oro linijų paslaugų sąžiningumas reikšmingai veikia vartotojų pasitenkinimą ir šis veiksnys turi didesnę poveikį nei suvokiama paslaugos kokybė. Tyrėjai taip pat pabrėžia, jog suvokiamas sąžiningumas ir vartotojų pasitenkinimas lemia didesnę pasitikėjimą oro linijų bendrovėmis. Akcentuojama, jog kainodaros sąžiningumas formuoja teigiamą vartotojų patirtį, didinant pasitenkinimą, o tai lemia didesnę pasitikėjimą teikiamomis paslaugomis.

Panašūs rezultatai atsispindi ir restoranų sektoriuje, Ahmed'o ir kt. (2023) tyrimas parodė, kad naudojantis restoranų paslaugomis, vartotojai jaučia didesnę pasitenkinimą, kai mano, jog kainodara yra sąžininga. Taip pat, nustatyta, jog restoranų paslaugos kokybė turi mažesnę poveikį vartotojų pasitenkinimui nei suvokiamas kainų sąžiningumas. Tyrimo rezultatai atskleidė, kad esant vartotojų požiūriu sąžiningoms kainoms, didėja vartotojų pasitenkinimas, o tai reikšmingai prisideda prie didėjančio lojalumo.

Alderighi ir kt. (2022) siekė identifikuoti veiksnius, lemiančius vartotojų suvokiamą kainų sąžiningumą viešbučių, ir turizmo sektoriuje bei nustatyti šio reiškinio poveikį vartotojų pasitenkinimui. Autoriai atskleidė, jog suvokiamam kainų sąžiningumui įtakos turi kainos dydis, kainos kintamumas, papildomi mokesčiai, viešbučio (pvz.: trijų, keturių ar penkių žvaigždučių) ir kambarių tipai. Nustatyta, kad viešbučiai, kurių kainodara suvokiama, kaip sąžininga, sulaukia didesnio vartotojų pasitenkinimo. Pabrėžiama, jog sąžininga kainodara turi reikšmingos įtakos vartotojų pasitenkinimui ir tai yra esminis verslo efektyvumo veiksnys.

Žinoma, jog vis sparčiau vartotojo ir paslaugą teikiančio žmogaus sąveiką keičia technologijos, tačiau suvokiamo sąžiningumo ir pasitenkinimo paslaugomis ryšys šioje tematikoje dar tik pradedamas nagrinėti. Mokslininkai Christ-Brendemühl ir Schaarschmidt'as (2022) nagrinėjo suvokiamo sąžiningumo poveikį vartotojų pasitenkinimui, analizuojant papildytos realybės (angl. *artificial reality, AR*) paslaugas apsiperkant internetu. Eksperimentinis tyrimas atskleidė, jog paskirstomojo, procedūrinio ir kainos sąžiningumo požiūriais vartotojai prasčiau vertino papildytos realybės paslaugas, lyginant su tradicinėmis, žmogiškųjų darbuotojų teikiamomis paslaugomis fizinėse apsipirkimo vietose. Nustatyta, jog suvokiamas papildytos realybės paslaugų nesąžiningumas prisidėjo prie prastesnės vartotojų patirties, mažesnio įsitraukimo, pasitenkinimo bei negatyvios komunikacijos iš lūpų į lūpas (angl. *word-of-mouth*). Autoriai įžvelgia mokslinės literatūros spragas, akcentuojant, jog vis dar mažai žinoma, kaip technologijos veikia vartotojų suvokiamą sąžiningumą ir su tuo susijusias reakcijas paslaugose.

Koh ir kt. (2025) analizavo, kokios DI pokalbių robotų savybės veikia vartotojų suvokiamą sąžiningumą, pasitikėjimą ir pasitenkinimą sutelktinio veikimo paslaugomis (angl. *crowdsourcing services*). Tyrimas atskleidė, jog DI pokalbių robotui suteikus žmogiškų savybių, reikšmingai didėja vartotojų suvokiamas sąžiningumas procedūrinio, paskirstomojo ir tarpasmeninio teisingumo požiūriais. Nustatyta, jog šios suvokiamo sąžiningumo dimensijos prisideda prie didėjančio vartotojų pasitikėjimo. Tuo tarpu pasitikėjimas DI pokalbių robotais veda link didesnio pasitenkinimo paslaugomis. Todėl, vartotojo ir DI sąveikoje suvoktas DI antropomorfizmas lemia numanomą technologijos sąžiningumą, o tai turi įtakos vartotojų pasitikėjimui ir pasitenkinimui paslaugomis.

Yang'as ir Lee'as (2024) siekė įvertinti algoritminio sąžiningumo vaidmens DI grįstose finansinėse paslaugose poveikį vartotojų pasitenkinimui ir ketinimams rekomenduoti. Tyrėjai nustatė, jog suvokiamą sąžiningumą lemia etinės DI savybės, tokios, kaip skaidrumas, atskaitomybė ir teisėtumas. Įrodyta, jog suvokiamas DI priimtų sprendimų sąžiningumas reikšmingai prisideda prie vartotojų pasitenkinimo paslaugomis, o tai turi įtakos vartotojų polinkiui rekomenduoti paslaugas kitiems. Todėl, autoriai akcentuoja, jog verslams būtina atsižvelgti į DI etines savybes, nes jos skatina suvokiamą sąžiningumą, o tai yra vienas iš pagrindinių veiksnių formuojant teigiamą vartotojų patirtį sąveikoje su technologijomis paslaugose.

Empiriniai duomenys rodo, jog suvokiamas sąžiningumas yra glaudžiai susijęs su vartotojų pasitenkinimu. Be to, vartotojų pasitenkinimas siejamas ir su kitomis teigiamomis atsako formomis – padidėjusiu pasitikėjimu, lojalumu bei ketinimais skleisti informaciją iš lūpų į lūpas. Mokslinė literatūra atskleidžia, jog vartotojų suvokiamam sąžiningumui įtakos gali turėti daugelis susijusių veiksnių, įskaitant paslaugų teikėjo tipą. Nepaisant to, vis dar trūksta žinių, kokį vaidmenį atlieka DI ir žmogaus sprendimų sąžiningumas vartotojų pasitenkinimui paslaugose.

Apibendrinant, galima teigti, jog vartotojų suvokiamas sąžiningumas turi didelę reikšmę vartotojų pasitenkinimui paslaugomis. Apžvelgus rinkodaros mokslo tyrimus, pastebima, jog egzistuoja spraga, analizuojant, kaip vartotojai suvokia DI ir žmogaus priimtų sprendimų sąžiningumą, ypač moralinėse paslaugų aplinkybėse. Todėl, atsižvelgiant į tyrimų išvadas kitose mokslo srityse, daroma prielaida, jog skirtingo paslaugos teikėjo tipo – žmogaus arba DI – priimtas sprendimas vartotojų požiūriu gali būti suvokiamas ir vertinamas nevienodai. Todėl, šiame projekte bus siekiama atsakyti į klausimą, kaip suvokiamas žmogaus arba DI sąžiningumas veikia vartotojų pasitenkinimą, moraliniu požiūriu svarbiose paslaugų aplinkybėse.

2.6. Konceptualus vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis modelis

Šioje darbo dalyje apibendrinami ankstesniuose poskyriuose aptarti mokslinės literatūros konstruktai, keliamos hipotezės bei pristatomas konceptualus modelis, nurodantis tyrimo kryptį bei ryšius tarp kintamųjų, kuriuos bus siekiama patikrinti empirinio tyrimo metu.

Vartotojų polinkis priskirti moralinį veiksnumą DI ir žmogaus sprendimams. Žinoma, kad vartotojai, moraliniu požiūriu svarbiose paslaugų teikimo aplinkybėse yra linkę moraliai svarstyti ir vertinti tarpusavio sąveikoje dalyvaujančius subjektus. Tačiau, rinkodaros mokslo srityje vis dar trūksta žinių, vertinant vartotojų polinkį priskirti moralinį veiksnumą DI ir žmogui. Esama mokslinė literatūra atskleidžia, jog moraliai svarbiose paslaugų aplinkybėse, suvokiamas subjekto moralinis veiksnumas gali reikšmingai prisidėti prie teigiamo arba neigiamo vartotojų atsako, pavyzdžiui didesnio / mažesnio pasitikėjimo paslauga arba paslaugos priimtumo (Wester, 2024).

Pastebima, jog verslai vis sparčiau diegia DI sistemas, siekiant automatizuoti paslaugų teikimo procesus (Statista, 2024), tokiu būdu keičiant įprastą, vartotojų ir žmoniškųjų darbuotojų sąveiką paslaugų sandoriuose (Pitardi et al., 2022). Atsižvelgiant į tai, jog paslaugose iš esmės nėra galimybės išvengti nesusipratimų, įskaitant moraliniu požiūriu svarbias aplinkybes, to negalima užtikrinti ir sąveikoje su DI sistemomis (Arikan et al., 2023). Tampa svarbu analizuoti, kuriam paslaugos teikimo subjektui, DI ar žmogui, vartotojai bus linkę priskirti didesnę moralinį veiksnumą. Siekiant suprasti kaip vartotojai vertina DI ir žmogų moraliai abejotinosiose paslaugų aplinkybėse, šį reiškinį geriausiai paaiškina suvokiamo moralinio veiksnumo (SMV) sąvoka. Suvokiamas moralinis veiksnumas grindžiamas vartotojų polinkiu priskirti moralines ir veiksnumo savybes stebimam subjektui. Siekiant įvertinti polinkį priskirti moralines savybes DI ir žmogui, remiamasi moralinių pagrindų teorija (MPT). Ši teorija nurodo, jog žmonės yra linkę vertinti kitų moralę, remiantis pagrindiniais moralės principais ir jų atitiktimi savoms moralinėms vertybėms. MPT apibrėžia tokius moralės pagrindus, kaip rūpestis / žala, sąžiningumas / apgavystė, lojalumas / ištikimybė, autoritetas / subversija, tyrumas / nuopuolis (Banks, 2019). Tuo tarpu vartotojų polinkis priskirti veiksnumo savybes grindžiamas proto teorija (PT). Ši teorija teigia, kad žmonės yra linkę priskirti psichines būsenas sau ir stebimiems subjektams, jas palyginti ir nuspėti galimą jų elgseną. Šiame magistro projekte remiamasi SMV, siekiant įvertinti vartotojų polinkį priskirti moralinį veiksnumą DI ir žmogaus sprendimams paslaugose.

Kitose mokslo srityse atlikti tyrimai rodo, jog moraliniu požiūriu abejotinosiose situacijose, žmonės yra linkę priskirti didesnę moralumą DI nei žmogui. DI priimti moraliai neteisingi sprendimai vertinami palankiau nei identiški žmogaus sprendimai (Wilson et al., 2022; Shank, 2019). Net ir moraliniu požiūriu sudėtingose situacijose, pavyzdžiui, diskriminacijos atveju, žmonės yra linkę manyti, jog

DI, dėl savo prigimties, negali pažeisti moralės normų, vedinas tam tikrų, tyčinių ketinimų, priešingai nei žmogus (Bigman et al., 2022). Didžioji dalis empirinių tyrimų atskleidžia, jog žmonės yra linkę priskirti didesnę veiksnumą žmogui, lyginant su DI. Šis dėsningumas yra aiškinamas suvokimu, jog DI gali būti pasitelkiamas tik techninių žinių reikalaujančioms užduotims, o moraliai abejotinosiose paslaugų situacijose, svarbus gebėjimas veikti savarankiškai, atsižvelgiant į išorinius aspektus (Arikan et al., 2023). Tačiau, šios išvados nėra vienareikšmės. Naujausi tyrimai rinkodaros mokslinėje literatūroje rodo, kad paslaugose, vartotojai yra linkę dehumanizuoti žmogiškuosius darbuotojus (Söderlund, 2024). Tuo tarpu DI suteikus proto savybių bei manipuliuojant jo elgseną, gali padidėti suvokiamas DI veiksnumas (Li et al., 2022).

Atsižvelgiant į mokslinių tyrimų išvadas apie polinkį priskirti didesnę moralumą DI nei žmogui, dėl suvokiamo DI negebėjimo ignoruoti moralės normų bei polinkio griežčiau vertinti moraliai abejotinus žmogaus sprendimus, tikimasi, jog didesnis moralumas bus priskirtas DI sprendimo atveju. Nepaisant esamų išvadų apie polinkį priskirti didesnę veiksnumą DI, šiame projekte remiamasi plačiau paplitusiomis mokslinėmis išvadomis, jog žmonės yra linkę priskirti didesnę veiksnumą žmogui (Arikan et al., 2023; Hwang ir Won, 2022; Thellman et al., 2022; Pitardi et al., 2022; Garvey et al., 2023). Šis požiūris grindžiamas vartotoju suvokiama didesne žmogaus atsakomybe bei didesniu suvokiamu protu. Taip pat, remiantis PT, vartotojai bus linkę priskirti didesnę veiksnumą žmogui, dėl suvokiamų žmogaus proto savybių. Todėl, formuluojamos pirmosios hipotezės:

H1a: DI priimtas sprendimas (palyginus su žmogaus sprendimu) sukels polinkį priskirti didesnę moralumą sprendimo priėmėjui;

H1b: Žmogaus priimtas sprendimas (palyginus su DI sprendimu) sukels polinkį priskirti didesnę veiksnumą sprendimo priėmėjui.

Sprendimo palankumo poveikis vartotojų polinkiui priskirti moralinį veiksnumą. Žinoma, jog sprendimo palankumas skatina polinkį priskirti teigiamus atributus sprendimo priėmėjui (Choung et al., 2023). Moksliniai tyrimai atskleidžia, jog suvokiamas priimto sprendimo palankumas teigiamai koreliuoja su vartotojų suvokiamu subjekto veiksnumu (Söderlund, 2024). Be to, sprendimo (ne)palankumas siejamas su suvokiamomis subjekto moralinėmis savybėmis. Empiriniai duomenys atskleidžia, jog paslaugose, priimto sprendimo (ne)palankumas, lemia vartotojų reakciją į paslaugos sandoryje dalyvaujančius subjektus, ypač nepageidaujamo rezultato atveju. Esant nepalankiam rezultatui, vartotojai yra linkę išsamiau analizuoti sprendimą lėmusias priežastis, nagrinėti išorinius veiksnus, įskaitant sprendimo priėmėjo savybes (Choi & Chao, 2024). Remiantis šiomis išvadomis, keliama prielaida:

H2a: Palankus sprendimas (palyginus su nepalankiu) sukels polinkį priskirti didesnę moralumą sprendimo priėmėjui;

H2b: Palankus sprendimas (palyginus su nepalankiu) sukels polinkį priskirti didesnę veiksnumą sprendimo priėmėjui.

Moderuojantis sprendimo (ne) palankumo poveikis vartotojų polinkiui priskirti moralinį veiksnumą DI ir žmogaus sprendimams. Vertinant, kaip vartotojų reakcija, (ne)palankaus sprendimo atveju kinta priklausomai nuo sprendimo priėmėjo tipo – DI arba žmogaus – pastebimos dvejopos mokslinių tyrimų išvados. Choi ir Chao (2024) tyrimas parodė, kad palankūs DI ir žmogaus sprendimai vertinami panašiai, tuo tarpu nepalankūs DI sprendimai vertinami palankiau nei žmogaus,

o tai lemia suvokiamas didesnis DI sąžiningumas. Tuo tarpu Yalcin ir kt. (2022) tyrimas parodė, jog vartotojai palankiau reaguoja į žmogaus pateiktą palankų sprendimą lyginant su DI, o nepalankaus sprendimo atveju abu subjektai vertinami vienodai. Esama mokslinė literatūra atskleidžia, jog nėra vieningų išvadų, kaip vartotojai reaguoja į žmogaus ir DI pateiktus identiškus (ne)palankius sprendimus.

Atsižvelgiant į tai, jog egzistuoja nevienareikšmės mokslinių tyrimų išvados, o rinkodaros mokslo srityje sprendimo (ne)palankumo tema, vertinant vartotojo sąveiką su skirtingais paslaugų teikėjų tipais, dar yra nauja, būtina atlikti tolimesnius tyrimus, siekiant nustatyti, kaip šis veiksnys pasireiškia moraliai abejotinosiose paslaugų aplinkybėse. Todėl, šiame magistro darbe remiamasi Choi ir Chao (2024) išvadomis, keliant prielaidą, jog priimto sprendimo palankumas teigiamai veiks suvokiamą moralumą, nepriklausomai nuo sprendimo priėmėjo tipo, o nepalankaus sprendimo atveju, DI bus priskiriamas didesnis veiksnumas:

H3a: Kai sprendimas yra palankus, nepriklausomai nuo to, ar sprendimą priėmė DI, ar žmogus, sprendimo priėmėjas bus suvokiamas, kaip labiau moralus;

H3b: Kai sprendimas yra nepalankus, DI sprendimų priėmėjas (palyginus su sprendimų priėmėju žmogumi) bus suvokiamas, kaip labiau veiksnus.

Vartotojų polinkio priskirti moralinį veiksnumą sprendimo priėmėjui įtaka suvokiamam sąžiningumui. Tyrimai rodo, jog vartotojai yra linkę vertinti paslaugos teikėjo priimto sprendimo teisingumą (Mathur & Sarin, 2020), o ši išvada ypač aktuali moraliniu požiūriu svarbiose paslaugų aplinkybėse. Suvokiamas sąžiningumas siejamas su priimto rezultato pagrįstumu, tinkamumu bei atitikimu moralės principams (Mathur & Sarin, 2020). Remiantis sąžiningumo teorija, vartotojai, siekdami įvertinti priimto sprendimo sąžiningumą, bus linkę kvestionuoti subjekto atskaitomybę, atsižvelgiant ir į moralės principų laikymąsi. Moralinėse paslaugų teikimo aplinkybėse, suvokiamo sąžiningumo ir vartotojų polinkio priskirti moralinį veiksnumą ryšys gali būti aiškinamas Colquitt'o ir kt. (2001) išskirtomis procedūrinio ir tarpasmeninio sąžiningumo dimensijomis. Paslaugose, procedūrinis sąžiningumas siejamas su subjekto gebėjimu kontroliuoti procesus, savarankiškai priimti sprendimus, koreguoti rezultatus bei pasižymėti kitomis savybėmis, kurios yra artimos suvokiamam veiksnumui (Weith & Matt, 2022). Tuo tarpu tarpasmeninis sąžiningumas siejamas su moraliai teisinga paslaugų subjekto elgsena, nes analizuojamas sąveikos teisingumas, įvertinant priimto sprendimo tinkamumą, empatiškumą ir mandagumą (Koh, 2025). Tikimasi, jog polinkis priskirti moralinį veiksnumą teigiamai veiks vartotojų suvokiamą sprendimo priėmėjo sąžiningumą, dėl numanomo gebėjimo veikti sąmoningai, vadovaujantis etikos ir moralės normomis. Todėl, šiame projekte bus siekiama paaiškinti sąsajas tarp vartotojų polinkio priskirti moralinį veiksnumą ir suvokiamo DI ir žmogaus sąžiningumo, keliant hipotezes:

H4a: Vartotojų polinkis priskirti moralumą sprendimo priėmėjui daro teigiamą poveikį suvokiamam sąžiningumui;

H4b: Vartotojų polinkis priskirti veiksnumą sprendimo priėmėjui daro teigiamą poveikį suvokiamam sąžiningumui.

Sprendimo priėmėjo įtaka vartotojų suvokiamam sąžiningumui. Esama mokslinė literatūra atskleidžia, jog vartotojai gali nevienodai vertinti žmogaus ir DI priimtų sprendimų sąžiningumą. Wu ir kt. (2024) atskleidė, kad vartotojų manymu, DI paslaugų robotai pasižymi didesniu sąžiningumu

nei žmonės. Tuo tarpu yra tyrimų, kurie nurodo priešingai, jog didesnis sąžiningumas priskiriamas žmogui lyginant su DI (Noble et al., 2021). Šios nevienareikšmės mokslinių tyrimų išvados grindžiamos daugeliu susijusių veiksnių, kaip priimto sprendimo kompleksiskumas, sudėtingumas, pasitikėjimas subjektu ar suvokiamas šališkumas. Be to, įtakos gali turėti ir sociodemografiniai rodikliai. Tačiau, didžioji dalis tyrimų atskleidžia, jog sąžiningumo požiūriu, dažniau pasitikima DI sprendimais. Ši išvada neretai lemiamą žmonių nuomone, jog DI negali turėti išankstinių nuostatų, savanaudiškų ketinimų ar pasižymėti emocionalumu, todėl yra sąžiningesnis nei žmogus (Choi & Chao, 2024). Atsižvelgiant į tai, jog ši tema mažai nagrinėta rinkodaros mokslo srityje, ypač kaip vartotojų polinkis priskirti moralinį veiksnumą veikia suvokiamą sąžiningumą, būtina atlikti tolimesnius tyrimus. Todėl, remiantis apžvelgtais moksliniais darbais, keliamos hipotezės:

H5a: Sprendimo priėmėjas moderuoja ryšį tarp polinkio priskirti moralumą ir sąžiningumo: kai sprendimą priima žmogus, ryšys tarp moralumo ir sąžiningumo bus stipresnis palyginus su situacija, kai sprendimą priima DI;

H5b: Sprendimo priėmėjas moderuoja ryšį tarp polinkio priskirti veiksnumą ir sąžiningumo: kai sprendimą priima DI, ryšys tarp veiksnumo ir sąžiningumo bus stipresnis palyginus su situacija, kai sprendimą priima žmogus.

Suvokiamo sprendimo priėmėjo sąžiningumo poveikis vartotojų pasitenkinimui paslauga. Suvokiamas priimto sprendimo sąžiningumas yra glaudžiai siejamas su vartotojų reakcija į paslaugos sandoryje dalyvaujančius subjektus bei į paslaugą plačiaja prasme. Žinoma, jog suvokiamas paslaugos sandorio nesąžiningumas veikia vartotojų atsaką ilgalaikėje ir trumpalaikėje perspektyvoje (Wu et al., 2024). Dauguma mokslinių tyrimų atskleidžia, jog suvokiamas sąžiningumas turi reikšmingą teigiamą poveikį vartotojų pasitenkinimui įvairiuose paslaugų kontekstuose (Koh et al., 2025; Yang & Lee, 2024; Wang et al., 2021; Alzoubi et al., 2020; Setiawan et al., 2020; Christ-Brendemühl & Schaarschmidt, 2022; Ahmed et al., 2023). Remiantis ankstesnėmis empirinių tyrimų išvadomis, keliami prielaidai:

H6: Suvokiamas sprendimo priėmėjo sąžiningumas teigiamai veikia vartotojų pasitenkinimą paslauga.

Moderuojantis sprendimo priėmėjo poveikis suvokiamo sąžiningumo ir pasitenkinimo paslauga ryšiui. Mokslinės literatūros apžvalga atskleidžia, jog vis dar trūksta žinių, kaip vartotojai suvokia identiškų DI ir žmogaus priimtų sprendimų sąžiningumą ir ar egzistuoja skirtumai vartotojų pasitenkinimo požiūriu. Naujausiuose tyrimuose, šie skirtumai nagrinėjami tarp kitų technologijų, kaip papildyta realybė (angl. *artificial reality*), ir žmonių. Christ-Brendemühl ir Schaarschmidt'as (2022) tyrimas parodė, jog paskirstomojo, procedūrinio ir kainos sąžiningumo atveju, vartotojai prasčiau vertino papildytos realybės paslaugas nei žmogiškųjų darbuotojų. Suvokiamas papildytos realybės sistemos nesąžiningumas turėjo įtakos mažesniai vartotojų pasitenkinimui. Atsižvelgiant į tai, daroma prielaida, jog suvokiamas DI nesąžiningumas, palyginti su žmogaus, turės didesnę įtaką vartotojų pasitenkinimui:

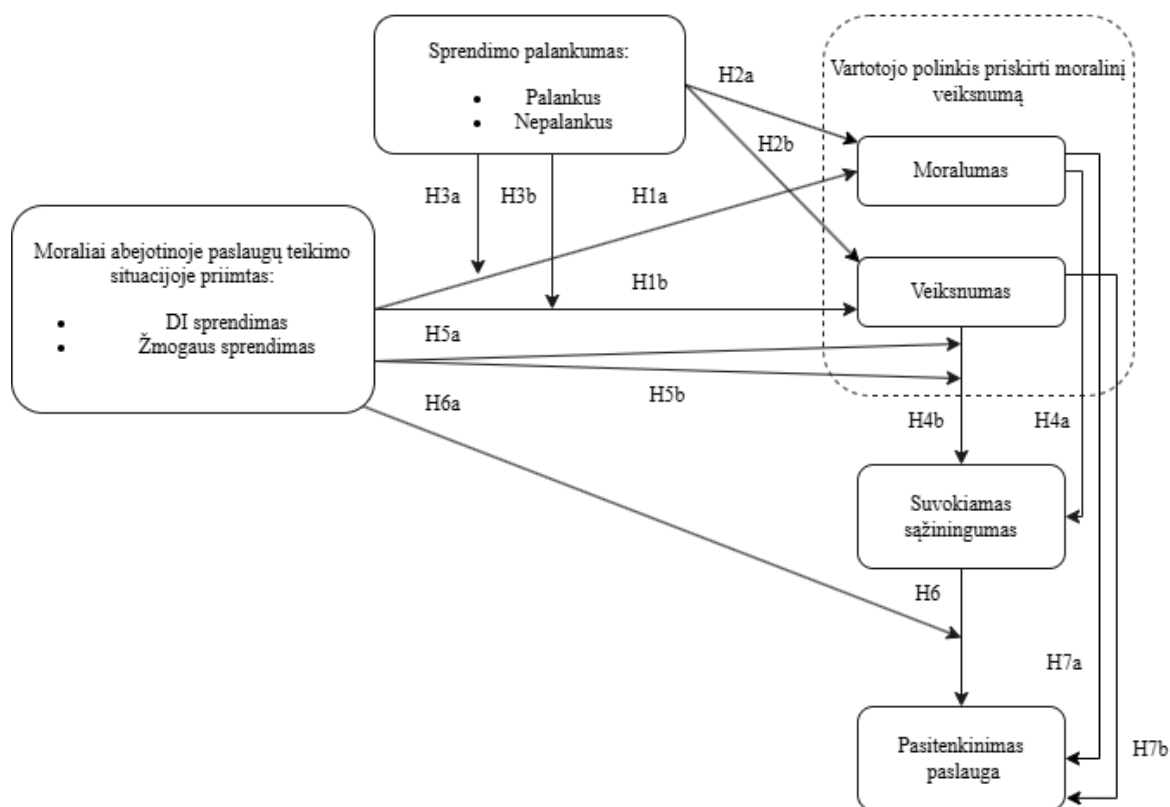
H6a: Sprendimo priėmėjas moderuoja ryšį tarp sąžiningumo ir pasitenkinimo taip, kad: kai sprendimą priima DI, ryšys tarp sąžiningumo ir pasitenkinimo bus stipresnis, palyginus su situacija, kai sprendimą priima žmogus.

Medijuojantis suvokiamo sąžiningumo poveikis tarp vartotojų polinkio priskirti moralinį veiksnumą ir pasitenkinimo. Naujausia literatūra atskleidžia, jog polinkis priskirti moralumą ir veiksnumą, lemia teigiamą vartotojų atsaką, kaip pasitikėjimą ir paslaugos priimtinumą (Zhao et al., 2025). Söderlund'as (2024) atskleidė, jog suvokiamas veiksnumas skatina teigiamus vartotojų atsiliepimus. Manoma, jog moralinis veiksnumas yra suvokiamo sąžiningumo sąlyga, nes sąžiningumas reikalauja gebėjimo priimti racionalius sprendimus, atsižvelgiant į moralės ir etikos normas. Įrodyta, kad suvokiamas sąžiningumas lemia vartotojų vertinimus ir atsaką paslaugose (Scholl-Grissemann et al., 2022), o šio reiškinio poveikis pasitenkinimui yra pripažintas daugelio autorių (Yang & Lee, 2024; Koh et al., 2025; Christ-Brendemühl & Schaarschmidt, 2022; Ahmed et al., 2023; Wang et al., 2021; Alzoubi et al., 2020; Setiawan et al., 2020;). Todėl, manoma, jog suvokiamas sąžiningumas atlieka ryšį tarp polinkio priskirti moralinį veiksnumą ir pasitenkinimo paslaugomis paaiškinančio mechanizmo vaidmenį:

H7a: Suvokiamas sąžiningumas medijuoja ryšį tarp moralumo ir pasitenkinimo taip, kad: moralumas daro teigiamą poveikį sąžiningumui, kuris savo ruožtu teigiamai veikia pasitenkinimą;

H7b: Suvokiamas sąžiningumas medijuoja ryšį tarp veiksnumo ir pasitenkinimo taip, kad: veiksnumas daro teigiamą poveikį sąžiningumui, kuris savo ruožtu teigiamai veikia pasitenkinimą.

Mokslinės literatūros apžvalga padėjo nustatyti ryšius tarp vartotojų polinkio priskirti moralinį veiksnumą DI arba žmogui, sprendimo palankumo, suvokiamo sąžiningumo ir vartotojų pasitenkinimo paslauga. Identifikuotas bei pagrįstas moderuojantis sprendimo palankumo poveikis polinkiui priskirti moralinį veiksnumą DI ir žmogaus sprendimams. Siekiant pagrįsti egzistuojančius ryšius tarp minėtų kintamųjų, sudarytas konceptualus modelis, kuris detalizuoja empirinio tyrimo kryptį (žr. **1 pav.**).



1 pav. Konceptualus vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikis pasitenkinimui paslaugomis modelis

3. Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis tyrimo metodologija

Šioje baigiamojo magistro projekto dalyje apibrėžiama empirinio tyrimo metodologija, išsikeliant tyrimo objektą, tikslą, uždavinius, nustatant tyrimo metodą ir metodiką, pristatoma tyrimo eiga bei duomenų analizės procedūros.

3.1. Empirinio tyrimo objektas, tikslas ir uždaviniai

Atsižvelgiant į mokslinių darbų ribotumą, nagrinėjant vartotojų polinkį priskirti moralinį veiksnumą DI ir žmogaus sprendimams, moraliai dviprasmiškose aplinkybėse, šio empirinio tyrimo kontekstas siejamas su moraliniu požiūriu abejotinu sprendimo priėmėjo atsaku į respondento užklausą. Šiame tyrime taikomi scenarijai adaptuoti pagal Obermeyer'io ir kt. (2019) tyrimą, kuris pristato realaus pasaulio pavyzdį – DI priimamų sprendimų šališkumą sveikatos priežiūros paslaugose. Mokslininkų tyrimas atskleidė, jog JAV sveikatos priežiūros sistemoje plačiai naudojamas DI algoritmas prioritetizavo europiečių kilmės pacientus, nepaisant to, jog afroamerikiečių kilmės pacientai susidūrė su didesnėmis sveikatos problemomis. DI rėmėsi istoriniais duomenimis, kurie nurodė, kad europiečių kilmės pacientai sveikatos priežiūros paslaugoms išleidžia daugiau. Remdamasis šiais duomenimis ir vertindamas pacientų ligų riziką, įtvirtino rasinį šališkumą (Obermeyer et al., 2019). Tai lėmė moraliniu požiūriu neigiamus DI sprendimų rezultatus, kuomet afroamerikiečių kilmės pacientams buvo skiriama perpus mažiau medicininių paslaugų. Šiame projekte, dėl ribotų tyrimo vykdymo aplinkybių, scenarijai pritaikyti lyties šališkumo problematikai.

Siekiant reprezentuoti pelno siekiantį verslo subjektą, nuspręsta DI šališkumo atvejį pritaikyti profesionalių grožio paslaugų industrijai. Remiantis statistiniais duomenimis, grožio salonų paslaugų rinka nuolat auga, o 2025 m. jos dydis sieks 128.6 mlrd. USD dolerių (Global Market Insights, 2024). Viena populiariausių profesionalių grožio paslaugų sritis siejama su plaukų priežiūra (Global Market Insights, 2024). Galima teigti, jog ši verslo sritis yra paplitusi plačioje vartotojų grupėje, įskaitant moteris ir vyrus. Empirinių tyrimų apžvalga atskleidžia, jog ankstesni moksliniai darbai panašioje tematikoje, dažniausiai nagrinėjo sveikatos (Sullivan & Wamba, 2022; Pitardi et al., 2022; Wester et al., 2022), restoranų (Ryoo et al., 2024; Sullivan & Wamba, 2022), viešbučių ir turizmo (Ryoo et al., 2024) paslaugas, o profesionalios grožio paslaugos dar menkai tyrinėtos. Toliau formuluojamas tyrimo objektas, tikslas ir empirinio tyrimo uždaviniai.

Tyrimo objektas – vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikis pasitenkinimui paslaugomis, profesionalių grožio paslaugų empiriniame kontekste.

Tyrimo tikslas – empiriškai pagrįsti vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikį pasitenkinimui paslaugomis, profesionalių grožio paslaugų atveju.

Empirinio tyrimo uždaviniai:

1. Apibrėžti sociodemografinės tyrimo imties charakteristikas;
2. Empiriškai patikrinti konceptualaus modelio konstrukto konceptualias struktūras;
3. Įvertinti sprendimų priėmėjo poveikį vartotojų polinkiui priskirti moralinį veiksnumą;
4. Įvertinti sprendimo palankumo poveikį vartotojų polinkiui priskirti moralinį veiksnumą;
5. Įvertinti kaip sprendimo palankumas veikia sprendimo priėmėjo poveikį polinkiui priskirti moralinį veiksnumą;

6. Nustatyti moralinio veiksnio poveikį suvokiamam sąžiningumui;
7. Nustatyti suvokiamo sąžiningumo poveikį pasitenkinimui profesionaliomis grožio paslaugomis;
8. Įvertinti, kaip sprendimų priėmėjas veikia ryšį tarp polinkio priskirti moralinį veiksnumą ir suvokiamo sąžiningumo;
9. Įvertinti, kaip sprendimų priėmėjas veikia ryšį tarp suvokiamo sąžiningumo ir pasitenkinimo profesionaliomis grožio paslaugomis;
10. Įvertinti netiesioginį polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikį vartotojų pasitenkinimui profesionaliomis grožio paslaugomis per suvokiamą sąžiningumą.

3.2. Empirinio tyrimo tipas, metodai ir tyrimo konstruktas

Tyrimo tipas. Siekiant empiriškai pagrįsti vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikį pasitenkinimui paslaugomis, nuspręsta taikyti kiekybinio tyrimo strategiją. Kiekybinio tyrimo pasirinkimas grindžiamas siekiu patvirtinti arba paneigti išsikeltas tyrimo hipotezes ir įgyti statistiniam tyrimui būdingus, apibendrintus vartotojų požiūrio bei elgsenos rezultatus (Habes ir kt., 2021).

Kiekybinio tyrimo strategijai įgyvendinti pasirinktas internetiniais scenarijais grįstas eksperimentas. Ladak ir kt. (2024) mokslinės literatūros apžvalga atskleidžia, jog nagrinėjant žmonių polinkį priskirti moralinį veiksnumą, dažniausiai taikomas kiekybinis tyrimas ir eksperimento dizainas, o duomenims rinkti neretai naudojama internetinė apklausa, grįsta scenarijais. Eksperimento dizaino tyrimas leidžia nagrinėti priežastinius ryšius tarp kintamųjų ir nustatyti: „[...] kaip kontroliuojami ir nekontroliuojami veiksniai yra susiję su proceso rezultatu (-ais).“ (Antony ir kt., 2020, p. 1159). Visa tai leidžia daryti išvadas apie priežasties ir pasekmės ryšius, itin aktualius ir vartotojų tyrimuose (Stoner, Felix ir Stadler, 2023). Taikant eksperimentinį tyrimo dizainą, siekiama geriau suprasti vartotojų požiūrį, reakciją, elgseną ir tai lemiančias priežastis. Internetiniais scenarijais grįsta apklausa bus siekiama nustatyti priežastinius ryšius.

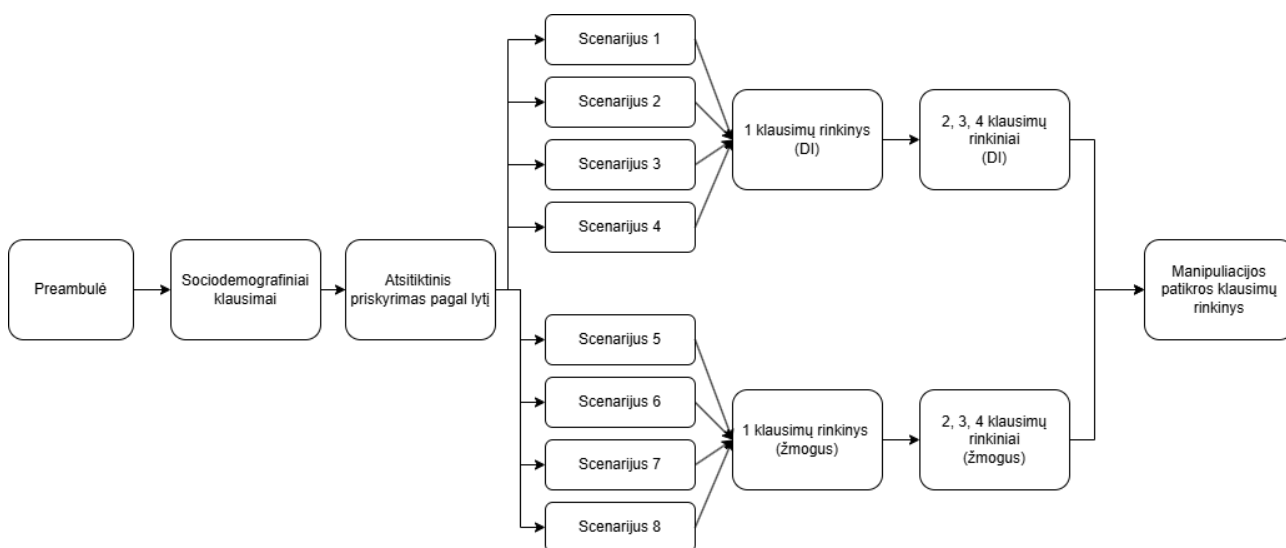
Duomenų rinkimo metodas. Atsižvelgiant į ankstesnių tyrimų metodologinę praktiką, tyrimo tikslą ir turimus išteklius, kiekybiniam tyrimui atlikti pasirinktas eksperimentinis tyrimo dizainas, taikant internetiniais scenarijais grįstą apklausos duomenų rinkimo metodą. Apklausos yra labiausiai paplitęs duomenų surinkimo būdas, leidžiantis rinkti duomenis apie vartotojų įsitikinimus, nuomones, nuostatas ir elgsenos ketinimus (Mohajan, 2020). Atliekant apklausas internetu, galima užtikrinti greitą, mažai išteklių reikalaujančią duomenų rinkimo procedūrą, leidžiančią pasiekti didesnę tyrimo imtį (Van Quaquebeke ir kt., 2022).

Internetinėje anketinėje apklausoje pateikiamas klausimynas, kurį sudaro aštuoni scenarijai ir tyrime naudojamų konstruktų matavimo skalės. Taikomas tarpgrupinis eksperimento dizainas, kuomet tyrimo dalyviai atsitiktinai priskiriami vienai iš eksperimentinių sąlygų. Respondentai skirstomi pagal 2 (sprendimo priėmėją: DI arba žmogus) $\times$ 2 (sprendimo palankumą: palankus arba nepalankus) tarpgrupinį dizainą. Siekiant užtikrinti, jog tyrimo dalyviai geriau įsijaustų į scenarijaus kontekstą, susijusį su lyties šališkumu, iš pradžių respondentai skirstomi pagal lytį į dvi grupes. Pradiniu respondentų suskirstymo pagal lytį prieš nukreipiant juos atsitiktiniu būdu į eksperimento sąlygas (scenarijus), siekta suaktyvinti moraliai abejotinos situacijos aplinkybes. Atsižvelgiant į minėtus aspektus ir taikant 2×2 tyrimo dizainą, iš viso sukurti aštuoni scenarijai. Siekiant užtikrinti galimybę skirstyti respondentų srautus atsitiktine tvarka, pasitelktas internetinių apklausų kūrimo įrankis

Qualtrics. Tyrime naudojamos ankstesniuose moksliniuose darbuose naudotos konstruktyvų matavimo skalės, kurios šiame projekte pritaikytos konkrečiam vykdomo tyrimo kontekstui.

Internetinė apklausa sudaryta anglų kalba ir pritaikyta plačiam geografiniam spektrui, reprezentuojant įvairias šalis. Atsižvelgiant į tyrimo tikslą, pobūdį bei turimus išteklius, imties atrankai nuspręsta taikyti netikimybinį patogųjį atrankos metodą.

Tyrimo instrumentas ir skalės. Internetinę apklausos anketą sudaro 11 blokų, apimančių preambulę, eksperimento grupių scenarijus, sociodemografinius, tyrimo konstruktyvų matavimo skalių, dėmesio (angl. *attention check*) ir manipuliacijos (angl. *manipulation check*) patikrinimo klausimus (žr. **2 pav.**). Anketos pradžioje, pateikiama informacija apie tyrimo tikslą bei etiką. Pirmuoju anketos klausimų bloku siekiama surinkti duomenis apie sociodemografines tiriamųjų charakteristikas. Iš kart po to, respondentai skirstomi pagal lytį, siekiant užtikrinti, jog tyrimo dalyviai labiau įsijaustų į moraliniu požiūriu abejotiną paslaugų teikimo situaciją, susijusią su lyties šališkumu. Respondentai, atsitiktine tvarka, kreipiami į vieną iš aštuonių (4x vyrų ir 4x moterų) tarpgrupinio eksperimento dizaino blokų, kuriuose pateikiami DI arba žmogaus, palankaus arba nepalankaus sprendimo scenarijai. Tuomet, respondentams pateikiami dėmesio patikros klausimai, kurie buvo adaptuoti iš ankstesnių tyrimų ir pritaikyti šio tyrimo kontekstui. Vėliau, pagal priskirtą scenarijų, respondentams pateikiami klausimai, pritaikyti scenarijuje aprašytam sprendimo priėmėjo tipui – DI arba žmogui. Kiekviename klausimų bloke, prašoma išreikšti savo nuomonę pritarant arba nepritarant pateiktiems teiginiams, reprezentuojantiems konstruktyvų matavimo skales. Polinkio priskirti moralinį veiksnumą konstruktas matuojamas dešimčia teiginių, suvokiamas sąžiningumas ir pasitenkinimas, kiekvienas, trimis teiginiais. Galiausiai, pateikiami manipuliacijos patikrinimo klausimai, kurie buvo adaptuoti esamam tyrimo kontekstui iš anksčiau atliktų empirinių tyrimų. Tyrimo dizainas ir anketos struktūra pateikiama 5 priede.



2 pav. Eksperimento dizainas ir anketos struktūra

Remiantis ankstesne literatūra (Wu et al., 2024; Choi & Chao, 2024), anketos pradžioje, pateikiamas tyrimo tikslas, pristatoma apklausos vykdymo eiga, pateikiama informacija apie tyrimo etiką. Siekiant užtikrinti respondentų objektyvumą, santraukoje neminimi tyrimui aktualūs konstruktai – vartotojų polinkis priskirti moralinį veiksnumą, sprendimo palankumo, suvokiamo sąžiningumo

vaidmuo ir pasitenkinimas paslauga. Tyrimo tikslas formuluojamas per siekį suprasti, kaip vartotojai vertina sprendimų priėmėjus, moraliai ginčytinose paslaugų situacijose.

Pirmoje anketos dalyje, renkama informacija apie respondentų sociodemografines charakteristikas, pateikiant 6 klausimus. Atsižvelgiant į tyrimo etiką bei tai, jog lyties pasirinkimo klausimo rezultatai turi įtakos tolimesnei apklausos eigai, respondentų paprašyta nurodyti biologinę lytį (angl. sex), pateikiant kategorinę skalę su dviem pasirinkimo variantais – vyro ir moters. Vėliau, paprašyta nurodyti išsilavinimo lygmenį, pateikiant ranginę skalę. Amžiui ir gyvenamajai vietai identifikuoti pateikiami atviri klausimai. Atsižvelgiant į ankstesnių studijų rezultatus, kurie nurodo, jog vartotojų reakcija į DI sprendimus gali priklausyti nuo technologinio raštingumo lygmens (Schoeffer et al., 2021a; Brailsford, 2024; Araujo et al., 2020), siekiama įvertinti respondentų patirtį, naudojantis DI paslaugomis. Šis klausimas adaptuotas pagal Yuan'o ir kt. (2024) tyrimą, pateikiant ranginę skalę. Galiausiai, siekiant įvertinti scenarijaus aktualumą abejoms respondentų lytims, remiamasi Koh ir kt. (2025) tyrimu, prašant nurodyti plaukų priežiūros paslaugų naudojimosi dažnį, pateikiant ranginę skalę.

Antrojoje anketos dalyje, respondentai skirstomi pagal lytį. Skirstymas pagal lytį grindžiamas siekiu sukurti sąlygas, leidžiančias geriau įsijausti į pateiktą scenarijų. Scenarijų kontekstas apima lyties šališkumo problematiką. Siekiant sustiprinti moraliai ginčytino sprendimo vartotojo atžvilgiu įspūdį, scenarijuose pateikiama situacija siejama su respondento lyčiai priešinga lytimi. Tuomet, atsitiktine tvarka, abiejų lyčių atstovams priskiriamas vienas iš 4 galimų scenarijų. Kiekviename scenarijuje respondentų prašoma įsijausti į grožio paslaugų teikimo situaciją, kurioje pateikus užklausą mobiliojoje grožio paslaugų programėlėje, inicijuojamas poreikis pasinaudoti plaukų priežiūros paslauga. Scenarijuose skiriasi sprendimo priėmėjas – DI arba žmogus, ir sprendimo palankumas (žr. 5 lentelė ir 6 priede):

- Pirmojo scenarijaus grupėje pateikiamas respondentui(-ei) palankus DI paslaugų agento sprendimas, apimantis teigiamą atsaką į respondento(-ės) pateiktą plaukų priežiūros paslaugų užklausą, pasinaudoti plaukų priežiūros paslauga;
- Antrojo scenarijaus grupėje pateikiamas respondentui(-ei) palankus žmogiškojo paslaugų darbuotojo sprendimas, apimantis teigiamą atsaką į respondento(-ės) pateiktą plaukų priežiūros paslaugų užklausą, pasinaudoti plaukų priežiūros paslauga;
- Trečiojo scenarijaus grupėje pateikiamas respondentui(-ei) nepalankus DI paslaugų agento sprendimas, apimantis neigiamą atsaką į respondento(-ės) pateiktą plaukų priežiūros paslaugų užklausą, pasinaudoti plaukų priežiūros paslauga;
- Ketvirtojo scenarijaus grupėje pateikiamas respondentui(-ei) nepalankus žmogiškojo paslaugų darbuotojo sprendimas, apimantis neigiamą atsaką į respondento(-ės) pateiktą plaukų priežiūros paslaugų užklausą, pasinaudoti plaukų priežiūros paslauga.

5 lentelė. 2x2 tarpgrupinio dizaino schema

Sprendimo palankumas	Sprendimo priėmėjas	
	DI	Žmogus
Palankus	DI paslaugų agento palankus sprendimas	Žmogiškojo paslaugų darbuotojo palankus sprendimas

Nepalankus	DI paslaugų agento nepalankus sprendimas	Žmogiškojo paslaugų darbuotojo nepalankus sprendimas
------------	--	--

Remiantis aprašytu DI rasių šališkumo atveju bei atsižvelgiant į vykdomo tyrimo galimybes, nuspręsta moraliai abejotiną paslaugų teikimo kontekstą sieti su lyties šališkumu. Moraliai abejotina paslaugų teikimo situacija siejama su sprendimo priėmėjo nesąžiningumu respondento(-ės) arba scenarijuje aprašyto, priešingos lyties kliento(-ės) atžvilgiu. Sprendimo priėmėjas – DI arba žmogus – skirstydamas klientų srautus, atsižvelgia į tokius kriterijus, kaip kliento(-ės) lytis, siekdamas didesnės naudos grožio salonui. Tyrimo scenarijai pateikti 6 priede.

Siekiant užtikrinti tyrimo duomenų kokybę, svarbu, jog respondentai dėmesingai įvertintų apklausoje pateikiamus klausimus. Anot Gummer'io ir kt. (2021), vykdant apklausas internetu, tampa itin sudėtinga užtikrinti tyrimo dalyvių dėmesį ir įsijautimą į hipotetinę situaciją, o tai gali neigiamai veikti tyrimo rezultatus. Šį reikalavimą padeda įgyvendinti dėmesio patikrinimo klausimai, nes jie leidžia įvertinti respondentų atidumą. Remiantis Wu ir kt. (2024) tyrimo praktika, tiek respondentų įsijautimo į situaciją sustiprinimui, tiek dėmesio patikrinimui („Apibūdinkite scenarijų bent dešimčia žodžių ir išsakykite savo mintis.“), naudojamas atviro tipo klausimas, kuriuo prašoma apibūdinti pateiktą scenarijų bent dešimčia žodžių. Tokio pobūdžio refleksija sustiprina tyrimo dalyvių įsijautimą į hipotetinę situaciją. Taip pat, dėmesio patikrai naudojami ir manipuliacijos klausimai, kuriais prašoma identifikuoti scenarijuje aprašytą sprendimo priėmėją bei priimto sprendimo rezultatą. Šioje anketos dalyje nuspręsta įtraukti ir scenarijaus realistiškumo įvertinimo klausimą, remiantis Wu ir kt. (2024) tyrimu. Respondentų prašoma atsakyti į klausimą („Kaip manote, kiek tikėtina, kad šis scenarijus įvyktų realiaame gyvenime?“), pateikiant 5 balų Likerto skalę, kur 1 = labai mažai tikėtina, o 5 = labai tikėtina. Šio klausimo įtraukimas leidžia patikrinti, ar tyrimo scenarijai respondentų buvo suvokiami kaip realistiški. Statistiniai šio klausimo atsakymų rezultatai prisideda prie tyrimo validumo užtikrinimo.

Trečioji anketos dalis skirta tyrimą reprezentuojančių konstrukto skalėms, siekiant įvertinti vartotojų polinkį priskirti moralinį veiksnumą, suvokiamą sąžiningumą ir pasitenkinimą paslauga. Pirmasis ir antrasis teiginių blokas sudarytas remiantis Banks (2019) dviejų dimensijų suvokiamo moralinio veiksnumo (angl. *perceived moral agency*) skale. Pirmasis teiginių blokas reprezentuoja pirmąją dimensiją – suvokiamą subjekto (DI arba žmogaus) moralę. Pastaroji matuojama 6 teiginiais. Antrasis teiginių blokas siejamas su antrąja dimensija – suvokiamu subjekto veiksnumu, kuri matuojama pateikiant 4 teiginius. Trečiuoju teiginių bloku siekiama įvertinti suvokiamą DI arba žmogaus sąžiningumą (angl. *perceived fairness*), pateikiant 3 teiginius. Suvokiamo sąžiningumo skalė adaptuota pagal Liu ir Sun'o (2024) bei Yang'o ir Lee (2024) sudarytas matavimo skales. Ketvirtasis teiginių blokas matuoja vartotojų pasitenkinimą, adaptuotą pagal Yang'o ir Lee (2024) matavimo skalę. Šia skalę sudaro 3 teiginiai.

Galiausiai, **ketvirtoje anketos dalyje**, siekiama įvertinti, ar tyrime naudotų nepriklausomų kintamųjų manipuliacija buvo interpretuojama taip, kaip tikėtasi, įtraukiant manipuliacijos patikrinimo klausimus. Respondentų pateikti atsakymai į šiuos klausimus leidžia įsitikinti, jog egzistuojantis poveikis priklausomiems kintamiesiems yra susijęs su manipuliacija (Viglia et al., 2021). Šiame projekte, manipuliacijos patikai atlikti remtasi vykdomam tyrimui aktualiu Choi ir Chao (2024) moksliniu darbu, pateikiant du klausimus, susijusius su sprendimo priėmėjo tipu ir sprendimo palankumu. Pirmuoju klausimu prašoma nurodyti, kas priėmė galutinį sprendimą – DI paslaugų agentas ar žmogiškasis paslaugų agentas. Antruoju klausimu prašoma identifikuoti koks buvo

galutinis sprendimas, vartotojui palankus – paslauga buvo suteikta anksčiau nei prašyta, ar vartotojui nepalankus – paslauga suteikta vėliau nei prašyta. Be to, siekiant įsitikinti, ar pateiktas scenarijus buvo suvoktas taip, kaip tikėtasi, įtrauktas klausimas („Kiek palankus galutinis sprendimas?“), prašant įvertinti priimto sprendimo palankumą 5 balų Likerto skalėje, kai 1 = labai nepalankus, o 5 = labai palankus. Remiantis Choi ir Chao (2024) empiriniu tyrimu, įtrauktas ir sprendimo priimtumo klausimas („Kiek priimtinas galutinis sprendimas?“), pateikiant 5 balų Likerto skalę, kai 1 = labai nepriimtinas, o 5 = labai priimtinas. Šiuo klausimu siekiama papildomai įvertinti sprendimo palankumą, iš priimtumo perspektyvos. Žemiau pateikiama, kaip buvo matuojami tyrimo konstruktai (žr. 6 lentelė).

6 lentelė. Tyrimo konstrukto matavimo skalės lietuvių kalba

Konstruktas		Skalės teiginiai	Autorius(-ai), metai
Suvokiamas moralinis veiksnumas (angl. <i>perceived moral agency</i>)	Moralumas	1. Šis [X] jaučia, kas yra gera ir bloga 2. Šis [X] gali apgalvoti, ar veiksmas yra moralus 3. Šis [X] gali jaustis įpareigotas elgtis moraliai 4. Šis [X] geba racionaliai vertinti gerį ir blogį 5. Šis [X] elgiasi pagal moralės taisykles 6. Šis [X] susilaikytų nuo veiksmų, kurie turėtų skaudžių pasekmių	Banks, 2019
	Veiksnumas	1. Šis [X] gali elgtis tik taip, kaip yra užprogramuotas elgtis. 2. Šio [X] veiksmai yra jo užprogramavimo rezultatas. 3. Šis [X] gali daryti tik tai, ką jam liepia žmonės. 4. Šis [X] niekada nedarytų nieko, ko nėra užprogramuotas daryti.	
Suvokiamas sąžiningumas		1. Manau, kad šis [X] su visais klientais elgiasi vienodai ir nediskriminuoja jų dėl asmeninių savybių, nesusijusių su plaukų priežiūros paslaugų veiksniais. 2. Tikiu, kad šis [X], priimdamas teisingus sprendimus, naudojasi patikimais ir nešališkais duomenų šaltiniais. 3. Manau, kad šis [X] priima nešališkus sprendimus be išankstinio nusistatymo ar favoritizmo.	Adaptuota pagal: Liu ir Sun (2024) Yang ir Lee (2024)
Vartotojų pasitenkinimas		1. Apskritai esu patenkintas [X] paslaugomis, kurias patyriau. 2. [X] paslaugos atitinka arba viršija mano lūkesčius sąžiningumo požiūriu. 3. Esu patenkintas [X] teikiamų paslaugų kokybe.	Adaptuota pagal: Yang ir Lee (2024)

Pastaba: X reikšmė kinta priklausomai nuo sprendimo priėmėjo tipo – DI paslaugų agento arba žmogiškojo paslaugų darbuotojo.

3.3. Empirinio tyrimo eigos ir duomenų analizės procedūros

Apžvelgus empirinio tyrimo tipą, metodus bei konstrukto matavimą, pristatoma tyrimo eiga ir duomenų analizės procedūros.

Tyrimo imtis. Siekiant apibrėžti tyrimo imtį, taikytas netikimybinis patogusis atrankos metodas. Šio metodo pasirinkimas grindžiamas tokia pačia ankstesnių empirinių tyrimų praktika, nagrinėjant DI ir žmogaus sprendimų vertinimus (Yalcin et al., 2022), vartotojų polinkį priskirti veiksnio savybes DI arba žmogui (Pitardi et al., 2022; Ryoo et al., 2024; Arıkan et al., 2023), suvokiamą DI

sąžiningumą ir šio reiškinių įtaką vartotojų ketinimams paslaugose (Garvey et al., 2023). Tyrimo imčių vertinimas priklauso nuo konkrečioje studijoje taikyto tyrimo dizaino, nes būtent dėl šios priežasties tyrimų imtis gali reikšmingai skirtis. Minėtuose empiriniuose tyrimuose taikyti 2x2, 3x2 arba 2x3 tarpgrupiniai tyrimo dizainai, todėl žemiau pateiktoje lentelėje, šalia kiekvieno tyrimo imties, pateiktas respondentų skaičius vienoje eksperimento dizaino grupėje (žr. **7 lentelė**).

7 lentelė. Nestatistinio lyginamojo DI ir žmogaus sprendimų rinkodaroje vertinimo susijusioje tematikoje imties dydžio nustatymas

Autorius(-iai), metai	Tyrimo dizainas	Tyrimo imtis	Respondentų skaičius vienoje eksperimento dizaino grupėje
Ryoo ir kt. (2024) – 1b tyrimas	2 x 2 tarpgrupinis dizainas	190 (prieš 220)	~48
Ryoo ir kt. (2024) – 2 tyrimas	3 x 2 tarpgrupinis dizainas	226 (prieš 250)	~38
Ryoo ir kt. (2024) – 3 tyrimas	3 x 2 tarpgrupinis dizainas	233 (prieš 268)	~39
Pitardi ir kt. (2022)	2 x 2 tarpgrupinis dizainas	166	~42
Pitardi ir kt. (2022)	2 x 2 tarpgrupinis dizainas	121	~30
Garvey ir kt. (2023) – 1a tyrimas	2 x 2 tarpgrupinis dizainas	174	~44
Garvey ir kt. (2023) – 1b tyrimas	2 x 2 tarpgrupinis dizainas	299	~75
Garvey ir kt. (2023) – 2 tyrimas	2 x 3 tarpgrupinis dizainas	698	~116
Garvey ir kt. (2023) – 3a tyrimas	2 x 2 tarpgrupinis dizainas	403	~101
Garvey ir kt. (2023) – 3b tyrimas	2 x 2 tarpgrupinis dizainas	400	100
Yalcin ir kt. (2022)	2 x 2 tarpgrupinis dizainas	303	~76
Arikan ir kt. (2023)	3 x 2 tarpgrupinis dizainas	270	45
Vidurkis:		290	63

Atlikus susijusių tyrimų apžvalgą nustatyta, jog tyrimų imties vidurkis yra 290, o vieną eksperimento dizaino grupę vidutiniškai sudaro 63 respondentai. Atsižvelgiant į šiame tyrime taikomą 2x2 dizainą, nestatistinis imties dydis šiame tyrime turėtų bent 251. Tuo tarpu vienoje eksperimento dizaino grupėje turėtų būti ne mažiau nei 63 respondentai.

Siekiant užtikrinti statistinį pagrįstumą, pasitelkiamas ir statistinis, A-priori imties dydžio nustatymo metodas. A-priori metodas buvo naudojamas Pavone ir kt. (2023), Arikan ir kt. (2023), Sullivan ir Wamba (2022) eksperimento dizaino tyrimuose. Šių metodų derinimas grindžiamas panašių tyrimų naudojamais imties dydžių nustatymo metodais. A-priori imties dydžio nustatymui pasitelkta G*Power 3.1.9.7_143 versijos skaičiuoklė. Atsižvelgiant į vykdomo tyrimo dizainą, taikytas dviejų faktorių (two-way ANOVA) imties skaičiavimo metodas. Siekiant užtikrinti statistinį patikimumą, pasirinktas efekto dydis (f) yra 0,25, klaidos tikimybė (α) yra 0,05, galios dydis ($1-\beta$) yra 0,95. Remiantis skaičiuotės pateiktais rezultatais, bendra tyrimui reikalinga imtis yra 210 (žr. **7 priede**).

Įvertinus dviejų imties skaičiavimo metodų rezultatus ir išvedus vidutinę imties dydžio reikšmę, nustatyta, jog šio tyrimo imtis turi būti ne mažesnė nei 231, o vienoje eksperimento dizaino grupėje turi būti ne mažiau nei 58 respondentai (žr. **8 lentelė**).

8 lentelė. Empirinio tyrimo imties dydžio nustatymas

Imties skaičiavimo metodas	Tyrimo imtis	Respondentų skaičius vienoje eksperimento dizaino grupėje
Nestatistinis lyginamasis	251	63
Statistinis A-priori	210	53
Vidurkis:	231	58

Anketos testavimas. Naujausiuose moksliniuose darbuose, įvairiuose tyrimų vykdymo etapuose, pradėti naudoti didieji multimodaliniai modeliai (angl. *large multimodal models*). Pasitelkiant tokias priemones, kaip generatyvinio dirbtinio intelekto sistemos (GDI), siekiama greitai ir ekonomiškai efektyviai imituoti bei prognozuoti žmogiškųjų tyrimo dalyvių atsaką, sukuriant sintetinius duomenų rinkinius. Empirinių rinkodaros ir vartotojų tyrimų apžvalga atskleidžia, jog sintetiniai duomenys gali būti itin veiksmingi atliekant bandomuosius tyrimus (Sarstedt et al., 2024). Panašiai teigia ir Yoo ir kt. (2025), pabrėždami, jog sintetinės imtys (angl. *silicon sample*) padeda interpretuoti tyrimo medžiagą prieš platesnę sklaidą. Be to, akcentuojama, jog šiame etape GDI modeliai gali suteikti įžvalgų apie tyrimo instrumento korekcijas.

Nepaisant to, jog GDI teikia išskirtines galimybes vykdant mokslinius tyrimus, egzistuoja ir su šių sistemų taikymu susijusios rizikos. Dirbtiniu intelektu paremti modeliai gali pateikti šališkus, realybės neatitinkančius rezultatus ir suformuoti klaidingas įžvalgas. Siekiant išvengti minėtų rizikų ir užtikrinti kuo didesnę duomenų patikimumą, itin svarbu pateikti tinkamai suformuluotą užklausą (angl. *prompt*). Yoo ir kt. (2025) teigimu, vienas efektyviausių užklausos formulavimo būdų yra minčių grandinės (angl. *chain of thought prompting*) metodas. Šis metodas atkartoja žmogiškąjį problemų sprendimo būdą, pateikiant sudėtingą informaciją laipsniškai ir formuluojant užduotis vieną po kitos. Tokiu būdu užtikrinama, jog GDI tinkamai interpretuos užklausą ir pateiks išsamiausius rezultatus.

Neretai, sintetiniuose tyrimuose pasitelkiamas generatyvus iš anksto apmokytas transformatorius (angl. *generative pre-trained transformer, GPT*). Anot Yoo ir kt. (2025), atliekant tyrimo instrumento testavimą, vienas efektyviausių GDI modelių yra *ChatGPT-4*. Mokslininkų teigimu, *ChatGPT-4* gali greitai sukurti dideles sintetines imtis ir efektyviai įgyvendinti eksperimentinius tyrimus, pateikiant susijusius klausimus kiekvienam sukurtam tyrimo dalyviui.

Atsižvelgiant į minėtus aspektus ir turimus ribotus tyrimo išteklius, šiame projekte atliktas išankstinis anketos skalių testavimas pasitelkiant *ChatGPT-4* modelį. Kuriant užklausą (angl. *prompt*), remtasi Yoo ir kt.(2025) tyrimo praktika, kuri nurodo, jog detalus tyrimo tikslo, siekiamų rezultatų bei instrumento aprašymas turi įtakos *ChatGPT-4* modelio sugeneruotų duomenų realistiškumui. Tyrėjų sugeneruota bandomojo tyrimo užklausa adaptuota pagal šio konkretaus tyrimo kontekstą. Atsižvelgiant į tyrimo pobūdį, sukurtos dvi identiškos užklausos, pritaikytos pagal respondentų lytį. Vienoje jų pateikiami moteriškos lyties scenarijai, prašant imituoti šios lyties respondentus, kitoje – vyriškai lyčiai būdingi scenarijai ir respondentai. Šis sprendimas grindžiamas siekiu pateikti kuo aiškesnes instrukcijas ir nesuklaidinti GDI modelio pateikiant itin kompleksinę užklausą. Užklausos

pradžioje pateikta informacija apie tyrimo konstruktus ir galimus ryšius tarp jų, vėliau, inicijuota užklausa sukurti po 200 respondentų, tuomet pateikti scenarijai, iš kart po to – tyrimo konstruktus reprezentuojantys klausimai. Po kiekvieno klausimo nurodytos aiškiai suformuluotos užklaustos, pageidaujant nurodyti tyrimo rezultatus. Gauti rezultatai padėjo įsitikinti, jog tyrimo manipuliacija yra veiksminga ir tyrimo instrumentas yra patikimas. Be to, įvertinus gautus duomenis, buvo paprašyta pateikti rekomendacijas, kurios padėtų sustiprinti scenarijuose aprašytą paslaugų teikimo situaciją taip, jog ji būtų labiau moraliai dviprasmiška ir glaudžiai sietusi su tyrimo konstruktais. Atsižvelgiant į sugeneruotas išvalgas, scenarijuje aprašyta moraliai abejotina situacija buvo pakoreguota, įtraukiant skubos ignoravimą ir objektyvumo stoką, nepateikiant paaiškinančios paslaugos teikimo nesėkmės priežasties. Pateiktos užklaustos, *ChatGPT-4* sintetinio tyrimo rezultatai, sugeneruotos rekomendacijos ir scenarijaus paragrafų pokytis, prieinami atvirojo mokslo sistemoje OSF: https://osf.io/kjd4r/?view_only=0f60752682b141f785ff5e46fc664832.

Faktinė tyrimo imtis. Pirminiais duomenimis anketą užpildė 409 respondentai. Atlikus dėmesio patikros duomenų analizę, nustatyta, jog dalis respondentų neatitiko šiame projekte numatytų reikalavimų, todėl galutinę tyrimo imtį sudarė 313 respondentai.

Tyrimo etika. Internetinės anketos preambulėje tyrimo dalyviai informuoti apie tyrimo tikslą, apklausos struktūrą, vykdymo eigą, duomenų apdorojimo procesus. Taip pat, laikantis tyrimų etikos principų, dalyviai informuoti, jog dalyvavimas šioje apklausoje yra savanoriškas, o pasitraukti iš tyrimo galima bet kuriame apklausos etape. Be to, akcentuota, jog ši internetinė apklausa yra anonimiška, gerbiant kiekvieno tyrimo dalyvio konfidencialumą.

Tyrimo vykdymas. Anketa buvo talpinama internetinėje apklausų kūrimo platformoje *Qualtrics* ir viešai prieinama nuo 2025 m. balandžio 6 dienos. Ši platforma yra patikimas ir plačiai naudojamas apklausų kūrimo įrankis, siūlantis platų programinį funkcionalumą. Atsižvelgiant į vykdomo tyrimo dizainą, *Qualtrics* pasirinkta dėl galimybės formuoti loginę apklausos vykdymo tvarką bei užtikrinti atsitiktinį respondentų srautų priskyrimą vienai iš 4 eksperimentinių sąlygų, po suskirsytymo į vyrų ir moterų grupes. Internetinė apklausa buvo viešinama asmeniniais komunikacijos kanalais bei socialiniuose tinkluose *LinkedIn*, *Instagram* ir *Facebook*. Duomenų rinkimas truko 14 dienų, apklausa deaktyvuota 2025 m. balandžio 20 dieną.

Tyrimo duomenų statistinis apdorojimas ir analizavimas. Statistinė duomenų analizė atlikta naudojant *IBM SPSS* programinę įrangą. Analizės metu taikyti aprašomosios, faktorinės, neparametrinių testų, ryšių tarp kintamųjų (koreliacijos), dispersijos, bei paprastosios tiesinės ir daugialypės tiesinės regresijos, statistinės duomenų analizės metodai. Duomenų patikimumui vertinti pasitelktas Kaiser'io-Mejer'io-Olkin'o imties adekvatumo matas ir Bartleto sferiškumo testas. Tiriamajai faktorinei analizei atlikta Direct Oblimin rotacija. Skalių patikimumui vertinti naudotas Kronbacho alfa (angl. *Cronbach alpha*) koeficientas. Nepriklausomų imčių t-testas taikytas manipuliacijos patikros analizei. Vienos imties t-testas scenarijų realistiškumui įvertinti. Tyrimo konstruktų raiškos priklausomybei nuo sociodemografinių duomenų taikyti Mann-Whitney ir Kruskal Wallis testai. Homogeniškumo patikrai atlikti naudotas Chi-kvadrato, Kruskal Wallis, vienos krypties tarpgrupinis ANOVA (angl. *one-way between groups Anova*) testai. Duomenų normalumas tikrintas Kolmogorovo-Smirnov (K-S) testu. Koreliacinei analizei taikytas Spearmano koreliacijos koeficientas. Dispersinei analizei pasitelkta dvipusė tarpgrupinė nepriklausomų imčių ANOVA (angl. *two-way between groups Anova*). Atlikta tiesinė paprastoji ir tiesinė daugialypė regresija. Regresinės

analizės moderavimo ir mediavimo efektai įvertinti naudojant *SPSS IBM* Hayes'o *PROCESS* v5.0 makrokomandos modelius Nr. 1 ir Nr. 4.

4. Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis empirinio tyrimo rezultatai ir diskusija

Šioje projekto dalyje pristatomi empirinio tyrimo rezultatai, pateikiama duomenų interpretacija, tikrinamos hipotezės. Be to, rezultatai apibendrinami tyrimo diskusijoje, apibrėžiami tyrimo ribotumai, tolimesnių tyrimų kryptys ir pasiūlymai.

4.1. Bendrosios tyrimo dalyvių sociodemografinės charakteristikos

Siekiant tinkamai interpretuoti rezultatus, pirmiausia analizuojamos vykdomam tyrimui aktualios sociodemografinės tyrimo dalyvių charakteristikos. Šioje darbo dalyje apžvelgiama respondentų lytis, išsilavinimo lygmuo, amžius, gyvenamoji vietovė, patirtis naudojantis DI paremtais sprendimais paslaugose bei naudojimosi plaukų priežiūros paslaugomis dažnis. Visi surinkti duomenys pateikti **8 priede**.

Tyrimo dalyviams iškelti dėmesio ir manipuliacijos patikros reikalavimai, užtikrinantys, jog respondentai vadovavosi apklausoje nurodytomis instrukcijomis ir atidžiai pildė klausimyną. Antroje anketos dalyje, prašyta apibūdinti scenarijuje pateiktą kontekstą, ketvirtoje, nurodyti scenarijuje aprašytą sprendimo priėmimą ir priimto sprendimo rezultatą. Žemiau pateiktoje lentelėje pristatomi duomenys, nurodantys respondentų pasiskirstymą pagal šių sąlygų atitikimą (žr. **9 lentelė**).

9 lentelė. Dėmesio ir manipuliacijos patikros duomenys (N=409)

Dėmesio ir manipuliacijos patikros rezultatai	Patikros klausimas	Teisingų atsakymų dažnis	Procentinė dalis
	Scenarijaus apibūdinimo klausimas	397	97,1%
	Sprendimo priėmimo klausimas	357	87,3%
	Sprendimo rezultato klausimas	349	85,3%
	Visų patikros klausimų suma:	313	76,5%

Gauti rezultatai atskleidžia, jog didžioji dalis teisingai atsakė į atskirus dėmesio ir manipuliacijos patikros klausimus. 97,1% respondentų tinkamai apibūdino pateiktą scenarijų, 87,3% nurodė sprendimo priėmimą (DI arba žmogų), 85,3% tinkamai nurodė priimto sprendimo rezultatą (paslauga suteikta anksčiau arba vėliau). Nepaisant to, nustatyta, jog į visus tris klausimus teisingai atsakė 76,5% respondentų. Todėl, į tolimesnę tyrimo rezultatų analizę įtraukiami tik tie respondentai, kurie teisingai atsakė į visus patikros klausimus (N=313). Detalesni rezultatai pateikiami **11 priede**.

Tyrimo duomenys atskleidžia, jog didžiąją dalį (61,3%) visų respondentų (N=313) sudarė moterys, tuo tarpu vyrų, kiek mažiau – 38,7%. Respondentų pasiskirstymas pagal išsilavinimą rodo, jog 59,7% tyrimo dalyvių yra įgiję bakalauro laipsnį, tuo tarpu mažiausia dalis, turi žemesnį nei vidurinį išsilavinimą (1,3%) (žr. **10 lentelė**).

10 lentelė. Demografinės tyrimo dalyvių charakteristikos (N=313)

Demografiniai rodikliai	Respondentų pasiskirstymas	Imtis (N)	Procentinė dalis
Lytis	Vyras	121	38,7%
	Moteris	192	61,3%

Išsilavinimas	Žemesnis nei vidurinis	4	1,3%
	Vidurinis arba lygiavertis atitikmuo	64	20,4%
	Bakalauro laipsnis arba lygiavertis atitikmuo	187	59,7%
	Magistro laipsnis arba lygiavertis atitikmuo	52	16,6%
	Daktaro laipsnis	6	1,9%

Atsižvelgiant į tai, jog tyrimas buvo vykdomas tarptautiniu mastu, respondentų buvo prašoma įvardinti šalį, kurioje šiuo metu reziduoja. Remiantis gautais duomenimis, respondentai atstovavo 27-nias šalis iš Europos, Azijos ir Amerikos (žr. **8 priede**). Siekiant apibendrinti gautus duomenis, išskirtos keturios pagrindinės šalys, pagal kurias pasiskirstė didžioji dalis respondentų (žr. **11 lentelė**). Rezultatai rodo, jog dauguma respondentų yra lietuviai, sudarydami 60,1% visų tyrimo dalyvių. 19,2% respondentų yra iš Jungtinės Karalystės, 2,9% iš Lenkijos, 1,6% atstovauja Prancūziją, o 16,3% kitas šalis, kaip Norvegija, Portugalija, Turkija ar Filipinai.

11 lentelė. Tyrimo dalyvių pasiskirstymas pagal šalį (N=313)

Šalis	Respondentų pasiskirstymas	Imtis (N)	Procentinė dalis
	Lietuva	188	60,1%
	Jungtinė Karalystė	60	19,2%
	Lenkija	9	2,9%
	Prancūzija	5	1,6%
	Kita	51	16,3%

Apžvelgus tyrimo dalyvių amžiaus duomenis, nustatyta, jog jauniausi respondentai buvo 18-likos metų, tuo tarpu vyriausi 55-erių. Vidutinis tyrimo dalyvių amžius buvo apie 27-nerius metus. Standartinio nuokrypio rodmeniu (6,90030%) atskleidžia, jog pagal amžių, respondentai buvo pasiskirstę nevienodai (žr. **12 lentelė**).

12 lentelė. Tyrimo dalyvių pasiskirstymas pagal amžių (N=313)

Amžius	Minimali vertė	Maksimali vertė	Vidurkis	Standartinis nuokrypis
	18	55	26,87	6,90030%

Siekiant suskirstyti tyrimo dalyvius į grupes, pasitelkta *SPSS Visual Binning* funkcija. Atsižvelgiant į tai, jog respondentai pasiskirstę netolygiai, grupės sudarytos pagal kvartilius, suskirstant tyrimo dalyvius į keturias amžiaus grupes (žr. **13 lentelė**). Gauti rezultatai atskleidžia, jog didžioji dalis (36,7%) respondentų buvo 18-23 m., tačiau didžiausia tyrimo dalyvių koncentracija buvo 24-25 m. amžiaus grupėje.

13 lentelė. Tyrimo dalyvių pasiskirstymas pagal amžiaus grupes (N=313)

Amžius	Respondentų pasiskirstymas	Imtis (N)	Procentinė dalis
	18-23 m.	115	36,7%
	24-25 m.	64	20,4%

	26-29 m.	68	21,7%
	30-55 m.	66	21,1%

Tyrimo dalyvių prašyta atskleisti savo naudojimosi DI paremtais sprendimais paslaugose patirtį. Gauti duomenys rodo, jog dauguma jų, 31,9%, tokiomis paslaugomis naudojasi nuo 1-nerių iki 2-ejų metų. Mažiausia dalis respondentų (10,2%) nurodė, jog patirties, naudojantis DI paslaugomis neturi (žr. **14 lentelė**).

14 lentelė. Respondentų pasiskirstymas pagal patirtį, naudojantis DI paremtais sprendimais paslaugose (N=313)

Patirtis, naudojantis DI paremtais sprendimais paslaugose	Respondentų pasiskirstymas	Imtis (N)	Procentinė dalis
	Jokia	32	10,2%
	Mažiau nei 1 metai	73	23,3%
	Nuo 1 iki 2 metų	100	31,9%
	Nuo 2 iki 3 metų	58	18,5%
	Daugiau nei 3 metai	50	16%

Siekiant įvertinti, kaip dažnai respondentai naudojasi tyrimo scenarijuje aprašytais plaukų priežiūros paslaugomis, prašyta nurodyti šių paslaugų vartojimo dažnį. Nustatyta, kad beveik ketvirtadalis, 32,6% tyrimo dalyvių, tokiomis paslaugomis naudojasi iki kelių kartų per metus. Tuo tarpu mažiausia dalis (2,9%) išreiškė, kad plaukų priežiūros paslaugomis naudojasi iki kelių kartų per savaitę (žr. **15 lentelė**).

15 lentelė. Respondentų pasiskirstymas pagal plaukų priežiūros paslaugų vartojimo dažnį (N=313)

Plaukų priežiūros paslaugų vartojimo dažnis	Respondentų pasiskirstymas	Imtis (N)	Procentinė dalis
	Beveik niekada	32	10,2%
	Iki kelių kartų per metus	182	32,6%
	Iki kelių kartų per pusmetį	87	27,8%
	Iki kelių kartų per mėnesį	83	26,5%
	Iki kelių kartų per savaitę	9	2,9%

Apžvelgus plaukų priežiūros paslaugų vartojimo tendencijas, siekiama nustatyti, kaip šiuo klausimu respondentai pasiskirstė pagal lytį. Šis siekis grindžiamas scenarijaus aktualumo įvertinimu vyrų ir moterų atžvilgiu. Gauti vidurkių palyginimo duomenys atskleidžia, jog vyrai tokiomis paslaugomis naudojasi dažniau, vidutiniškai nuo kelių kartų per pusmetį iki kelių kartų per mėnesį. Tuo tarpu moterys kiek rečiau, nuo kelių kartų per metus iki kelių kartų per pusmetį (žr. **16 lentelė**).

16 lentelė. Respondentų pasiskirstymas pagal lytį, naudojantis plaukų priežiūros paslaugomis (N=313)

Plaukų priežiūros paslaugų vartojimas pagal lytį	Respondentų pasiskirstymas	Vidurkis	Imtis (N)	Standartinis nuokrypis
	Vyrai	3,2975	121	,94556
	Moterys	2,4740	192	,96502

4.2. Tyrimo instrumento metodologinės kokybės įvertinimas

Tolimesnei statistinei duomenų analizei atlikti, svarbu įvertinti turimų duomenų tinkamumą. Tyrimo konstrukto skalių tinkamumui vertinti taikytas tiriamosios faktorinės analizės metodas (žr. **9 priede**). Skalių faktoriškumas vertintas remiantis Kaiser'io-Mejer'io-Olkin'o (KMO) imties adekvatumo matu ir Bartleto sferiškumo testo p reikšme (žr. **17 lentelė**).

17 lentelė. Tiriamų konstrukto skalių tinkamumo vertinimas (N=313)

Skalės pavadinimas	Teiginių skaičius	KMO imties adekvatumo matas	Bartleto sferiškumo kriterijaus p reikšmė
Polinkio priskirti moralinį veiksnumą skalė (moralumas)	6	0,915	<0,001
Polinkio priskirti moralinį veiksnumą skalė (veiksnumas)	4		
Vartotojų suvokiamo sąžiningumo skalė	3		
Vartotojų pasitenkinimo skalė	3		

KMO leidžia įvertinti koreliacijos stiprumą tarp kintamųjų, o faktorinę analizę prasminga atlikti tuomet, kai indekso reikšmė svyruoja nuo 0,6 iki 1,0 (Pallant, 2010). Bartleto sferiškumo reikšmė turi būti statistiškai reikšminga ($p < 0,05$), jog būtų patvirtintas faktorinės analizės naudingumas (Pallant, 2010). Vertinant šio tyrimo duomenų tinkamumą, KMO (0,915) ir Bartleto testo ($p < 0,001$) reikšmės nurodo, jog egzistuoja stipri reikšminga koreliacija tarp tiriamų konstrukto skalių.

Įvertinus duomenų tinkamumą, faktorinei analizei taikytas *Direct Oblimin* rotacijos metodas. Šio metodo pasirinkimas grindžiamas egzistuojančia stipria tyrimo konstrukto skalių koreliacija (Alordiah & Chenube, 2023). Taikant minėtą duomenų pasukimo metodą, išskirti faktoriai ir jų svoriai atskleidžia, jog kintamieji atitiko konceptualią struktūrą (žr. **18 lentelė**).

18 lentelė. Tiriamų konstrukto skalių faktorinės analizės rezultatai (N=313)

Skale matuojamas konstruktas	Skalės teiginiai	Išskirti faktoriai ir jų svoriai			
		1*	2*	3*	4*
Polinkis priskirti moralinį veiksnumą (moralumas)	Šis [X] jaučia, kas yra gera ir bloga	0,871			
	Šis [X] gali apgalvoti, ar veiksmas yra moralus	0,854			
	Šis [X] gali jaustis įpareigotas elgtis moraliai	0,816			
	Šis [X] geba racionaliai vertinti gerį ir blogį	0,812			
	Šis [X] elgiasi pagal moralės taisykles	0,735			
	Šis [X] susilaikytų nuo veiksmų, kurie turėtų skaudžių pasekmių	0,727			
Polinkis priskirti moralinį veiksnumą (veiksnumas)	Šis [X] gali elgtis tik taip, kaip yra užprogramuotas elgtis.		0,889		
	Šio [X] veiksmas yra jo užprogramavimo rezultatas.		0,873		
	Šis [X] gali daryti tik tai, ką jam liepia žmonės.		0,850		
	Šis [X] niekada nedarytų nieko, ko nėra užprogramuotas daryti.		0,842		

Suvokiamas sąžiningumas	Manau, kad šis [X] su visais klientais elgiasi vienodai ir nediskriminuoja jų dėl asmeninių savybių, nesusijusių su plaukų priežiūros paslaugų veiksniais.			-0,878	
	Tikiu, kad šis [X] naudoja patikimus ir nešališkus duomenų šaltinius teisingiems sprendimams priimti.			-0,868	
	Tikiu, kad šis [X] priima nešališkus sprendimus be išankstinio nusistatymo ar palankumo.			-0,850	
Pasitenkinimas	Apskritai esu patenkintas šio [X] suteiktomis paslaugomis.				0,915
	Šio [X] paslaugos atitinka arba viršija mano lūkesčius sąžiningumo požiūriu.				0,880
	Esu patenkintas šio [X] teikiamų paslaugų kokybe.				0,743

Polinkio priskirti moralinį veiksnumą konstrukto skalė skirstoma į dvi dimensijas – moralumo ir veiksnumo. Moralumo skalė matuojama 6-iais teiginiais. Visi skalės teiginiai priklauso pirmajam faktoriui. Veiksnumo skalės 4-eri teiginiai priklauso antrajam faktoriui. Suvokiamo sąžiningumo skalė matuojama 3-jais teiginiais, priklausydami trečiajam faktoriui. Pasitenkinimo skalės 3-eji teiginiai priklauso ketvirtajam faktoriui. Visi teiginiai paaiškina 78,55% dispersijos kintamuosiuose (žr. **9 priede**), todėl yra tinkami naudoti tolesnėse analizėse.

Siekiant įvertinti kiekvienos skalės patikimumą, nustatytas Kronbacho alfa (angl. *Cronbach alpha*) koeficientas (žr. **19 lentelė** ir **9 priede**).

19 lentelė. Tiriamų konstrukty skalių patikimumo vertinimas (N=313)

Skalės pavadinimas	Dimensija	Teiginių skaičius	Kronbacho alfa koeficientas
Polinkio priskirti moralinį veiksnumą skalė	Moralumas	6	0,934
	Veiksnumas	4	0,905
Suvokiamo sąžiningumo skalė		3	0,907
Pasitenkinimo skalė		3	0,952

Pallant (2010) nurodo, jog vidinis skalių nuoseklumas yra labai geras, kai α koeficiento reikšmė yra didesnė nei 0,8. Gauti rezultatai atskleidžia, jog visų skalių α reikšmės yra didesnės nei 0,9, tokiu būdu grindžiant tyrimo skalių patikimumą.

Remiantis faktorinės analizės rezultatais, nustatyta, jog tyrimo konstrukty struktūra atitiko teoriškai prognozuotą kintamųjų rinkinio struktūrą, o aukštas skalių patikimumas leidžia atlikti tolesnes statistines analizes ir patikrinti iškeltas hipotezes.

4.3. Polinkio priskirti moralinį veiksnumą, suvokiamo sąžiningumo ir pasitenkinimo kintamųjų charakteristikos

Prieš tyrimo hipotezių patikrą, svarbu apibūdinti tiriamų kintamųjų charakteristikas. Aprašomoji statistika padeda įvertinti parametrinių metodų prielaidas, tokiu būdu užtikrinant tolesnių analizių patikimumą (Pallant, 2010). Aprašomojoje tyrimo kintamųjų analizėje nurodomos minimalios, maksimalios, vidutinės reikšmės bei standartinis nuokrypis (žr. **20 lentelė** ir **10 priede**). Siekiant eliminuoti galimas klaidas, nustatant ryšius tarp kintamųjų, sukurti suminiai skalių vidurkių kintamieji.

20 lentelė. Tyrimo kintamųjų charakteristikos (N=313)

Skalės pavadinimas	Dimensija	Minimali reikšmė	Maksimali reikšmė	Vidurkis	Standartinis nuokrypis
Polinkio priskirti moralinį veiksnumą skalė	Moralumas	1	5	2,3600	1,04246
	Veiksnumas	1	5	3,5152	1,03522
Suvokiamo sąžiningumo skalė		1	5	2,2726	1,14395
Pasitenkinimo skalė		1	5	2,3600	1,33747

Vertinant polinkį priskirti moralinį veiksnumą, pastebima, jog moralumo dimensijos teiginiai vertinami labiau neigiamai (2,3600), o tai rodo, jog respondentai sprendimo priėmėjams buvo linkę priskirti mažesnę moralumą. Tačiau, standartinio nuokrypio reikšmė (1,04246) atskleidžia, jog respondentų atsakymai nebuvo vienareikšmiai. Tuo tarpu vertindami veiksnumo skalės teiginius, tyrimo dalyviai išreiškė didesnę pritarimą. Priskiriamo veiksnumo reikšmė artima 4, kuri pateiktame klausimyne atitiko dalinį sutikimą, jog sprendimo priėmėjai yra priklausomi nuo kitų (mažiau veiksnūs). Standartinis nuokrypis šioje skalėje yra mažiausias (1,03522), tačiau taip pat rodantis respondentų nuomonių išsiskyrimą. Polinkio priskirti moralinį veiksnumą skalių minimali reikšmė – 1, maksimali – 5.

Suvokiamo sąžiningumo vertinimas artimas 2, kuris nurodo, jog respondentai labiau nesutinka su pateiktais teiginiais, o standartinis nuokrypis (1,14395) rodo respondentų požiūrio skirtumus. Rezultatai suponuoja, jog pateiktas scenarijus buvo suvokiamas, kaip nesąžiningas. Minimali suvokiamo sąžiningumo skalės reikšmė yra 1, maksimali 5.

Respondentai dažniau išreiškė nepasitenkinimą, pasirinkdami 2-ą skalės reikšmę, kuri atitinka dalinį nepritarimą. Pasitenkinimo skalės standartinis nuokrypis (1,33747) yra didžiausias iš visų kintamųjų, todėl galima teigti, jog šiuo klausimu respondentų nuomonės išsiskyrė labiausiai. Minimali skalės reikšmė lygi 1, maksimali 5.

Dėmesio, manipuliacijos patikros ir scenarijų realistiškumo analizė. Kaip minėta anksčiau, dalis manipuliacijos patikros klausimų buvo naudoti tyrimo dalyvių dėmesio patikrinimui. Tyrimo imtyje liko tik tie respondentai, kurie teisingai atsakė į žemiau nurodytus klausimus, prašant nurodyti scenarijuje aprašytą sprendimo priėmimą ir priimto sprendimo rezultatą. Visą tai patvirtina gauti manipuliacijos patikros rezultatai (žr. **21 lentelė** ir **11 priede**), pavaizduoti naudojantis SPSS kryžminės lentelės (angl. *crosstab*) funkcija.

21 lentelė. Sprendimo priėmėjo ir sprendimo rezultato manipuliacijos patikros rezultatai (N=313)

Manipuliacija	Scenarijus	N	Teisingai atsakė, %	Neteisingai atsakė, %	Atsakymų suma, %
Sprendimo priėmėjo identifikaciją matuojantis kintamasis	DI	163	52,1%	0,0%	100%
	Žmogus	150	47,9%	0,0%	
Priimto sprendimo identifikaciją matuojantis kintamasis	Palankus	148	47,3%	0,0%	100%
	Nepalankus	165	52,7%	0,0%	

Siekiant įsitikinti, jog respondentai interpretavo tyrimo scenarijus taip, kaip numatyta, atlikta manipuliacijos patikros analizė, įvertinant nepriklausomų imčių t-testo reikšmes (žr. **22 lentelė**). Duomenys pateikti ir **11 priede**.

22 lentelė. Sprendimo palankumo ir sprendimo priimtumo vidurkių palyginimas (N=313)

Manipuliacija	Scenarijus	N	Vidurkis	Standartinis nuokrypis	Standartinė paklaida	Leveno testas – <i>p</i> reikšmė	Reikšmė <i>p</i> vidurkių skirtumams
Sprendimo palankumą matuojantis kintamasis	Palankus	148	4,36	0,801	0,066	<0,001	<0,001
	Nepalankus	165	1,33	0,576	0,045		<0,001
Sprendimo priimtumą matuojantis kintamasis	Palankus	148	3,65	1,194	0,098	<0,001	<0,001
	Nepalankus	165	1,44	0,710	0,055		<0,001

Nepriklausomų imčių t-testo rezultatai patvirtina tyrime naudotų manipuliacijų veiksmingumą. Sprendimo palankumo manipuliacijos vidurkių skirtumo reikšmė ($p < 0,001$) grindžia reikšmingą statistinį skirtumą tarp dviejų grupių (Pallant, 2010). Papildomai įtraukas sprendimo priimtumo klausimas dar kartą patvirtino, jog manipuliacijos turėjo numatytą poveikį ($p < 0,001$). Statistinis reikšmingumas atskleidžia, jog respondentai scenarijus interpretavo taip, kaip tikėtasi ir gebėjo identifikuoti pateiktų scenarijų palankumą ir nepalankumą.

Toliau atlikta scenarijaus realistiškumo patikra, kuri leidžia įvertinti tyrimo dalyvių požiūrį į pateiktų scenarijų kontekstą. Rezultatams atskleisti, naudotas vienos imties t-testas (žr. **23 lentelė** ir **11 priede**).

23 lentelė. Scenarijaus realistiškumo analizė (N=313)

Skalės pavadinimas	Vidurkis	Statistika (t)	Laisvės laipsniai (df)	<i>p</i> reikšmė
Scenarijaus realizmo kintamasis	3,27	4,175	312	<0,001

Scenarijaus realistiškumas patvirtinamas, kai vidurkis yra statistiškai reikšmingai didesnis už vidutinę skalės reikšmę (Wu et al., 2024). Remiantis gautais vienos imties t-testo duomenimis, scenarijaus realizmas įvertintas daugiau už vidurkį (3). Statistiškai reikšmingas skirtumas ($p < 0,001$) patvirtina, jog respondentų nuomone, aprašyti scenarijai gali įvykti realiame gyvenime. Šie duomenys leidžia daryti prielaidą, jog scenarijaus realistiškumas galėjo teigiamai veikti tyrimo dalyvių įsitraukimą ir pateiktų duomenų tikslumą.

Siekiant geriau interpretuoti tyrimo rezultatus, svarbu įvertinti tyrimo konstrukto raiškos priklausomybę nuo sociodemografinių tiriamųjų duomenų.

Tyrimo konstrukčių raiškos priklausomybės nuo sociodemografinių tiriamųjų charakteristikų analizė. Analizei atlikti taikyti Mann-Whitney ir Kruskal Wallis testai (žr. **12 priede**). Mann-Whitney testas taikytas dviejų nepriklausomų grupių kintamųjų skirtumams įvertinti. Tuo tarpu Kruskal Wallis testu analizuoti daugiau nei trijų grupių kintamieji – išsilavinimo ir patirties, naudojantis DI paremtomis paslaugomis. Abejų testų statistinis reikšmingumas patvirtinamas, kai p reikšmė yra mažesnė nei 0,05 (Pallant, 2010).

Siekiant nustatyti, ar egzistavo skirtumai tarp tyrimo dalyvių atsakymų priklausomai nuo lyties, taikytas neparametrinis Mann-Whitney testas (žr. **24 lentelė**).

24 lentelė. Tyrimo kintamųjų raiškos priklausomybė nuo lyties (N=313)

Skalės pavadinimas	p reikšmė	Lytis	Vidurkio rangas
Polinkio priskirti moralinį veiksnumą skalė (moralumas)	0,040	Vyras	170,13
		Moteris	148,72
Polinkio priskirti moralinį veiksnumą skalė (veiksnumas)	0,545	Vyras	160,88
		Moteris	154,55
Vartotojų suvokiamo sąžiningumo skalė	0,397	Vyras	162,37
		Moteris	153,62
Vartotojų pasitenkinimo skalė	0,135	Vyras	166,37
		Moteris	151,10

Analizuojant gautus rezultatus, pastebima, jog moterys ir vyrai sprendimo priėmėjų moralumą vertino nevienodai. Statistiškai reikšmingas ryšys ($p = 0,040$) grindžia egzistuojančius moralumo vertinimo skirtumus, o vidurkio rangai patikslina, jog moterys (170,13) buvo linkusios priskirti didesnę moralumą, lyginant su vyrais (148,72). Skirtumai tarp polinkio priskirti veiksnumą ($p = 0,545$), suvokiamo sąžiningumo ($p = 0,397$) ir vartotojų pasitenkinimo ($p = 0,135$) nebuvo pastebėti, nes visų šių analizuojamų kintamųjų p reikšmės yra didesnės nei 0,05.

Toliau siekta išsiaiškinti, ar pateikti atsakymai skyrėsi tarp skirtingų respondentų grupių pagal išsilavinimą. Analizei atlikti taikytas neparametrinis Kruskal Wallis testas. Šios tyrimo kintamųjų raiškos priklausomybės patikra grindžiama literatūroje pateikiamomis išvadomis, jog sprendimo priėmėjo vertinimas gali priklausyti nuo išsilavinimo lygmens, ypač suvokiamo sąžiningumo požiūriu (Grgić-Hlača et al., 2022). Kruskal Wallis testo rezultatai rodo, jog analizuojamuose kintamuosiuose nėra statistiškai reikšmingų skirtumų tarp skirtingų išsilavinimo grupių (žr. **12 priede**).

Galiausiai, analizuoti skirtumai, priklausomai nuo respondentų patirties, naudojantis DI paremtomis paslaugomis, taikant Kruskal Wallis testą. Prielaida, jog tyrimo dalyvių atsakymai gali skirtis pagal turimą patirtį, naudojantis DI sprendimais paslaugose, kelta remiantis empirinėmis išvadomis

(Brailsford, 2024; Araujo et al., 2020). Gauti rezultatai atskleidžia, jog reikšmingumas pagal minėtą sociodemografinę charakteristiką, tiriamuose kintamuosiuose neegzistuoja (žr. **12 priede**).

Toliau, tikrinama ar eksperimento grupės yra homogeniškos pagal sociodemografinės charakteristikas, įtrauktas į tyrimo instrumentą (žr. **13 priede**).

Eksperimento grupių homogeniškumo patikros rezultatai. Pirmiausia, patikra atlikta tarp dviejų grupių, kategorinio lyties kintamojo, naudojant Chi-kvadrato testą. Testo rezultatai atskleidžia, kad $p=0,996$, o tai rodo, kad vyrų ir moterų proporcijos, skirtinguose scenarijuose, statistiškai reikšmingai nesiskiria. Todėl, galima patvirtinti, jog **grupės pagal lytį yra homogeniškos**.

Rangų išsilavinimo kintamojo homogeniškumo patikrai taikytas Kruskal Wallis testas. Stebimas vidurkių rangų panašumas, o p reikšmė indikuoja, jog išsilavinimo vidurkių rangai, skirtingose eksperimento grupėse, statistiškai reikšmingai nesiskiria ($p = 0,401$). Remiantis šiuo požymiu galima teigti, jog **grupės pagal išsilavinimą išlieka homogeniškos**.

Kadangi amžius matuotas santykinė skale, taikytas vienos krypties tarpgrupinis ANOVA (angl. *one-way between groups Anova*) testas. Leveno testo p reikšmė patvirtina dispersijų lygybės prielaidą ($p=0,849$), ANOVA testas atskleidžia, jog statistiškai reikšmingų skirtumų pagal amžių, skirtingose eksperimento grupėse, neaptikta ($p = 0,996$). Tyrimo dalyvių amžiaus vidurkis atskirose eksperimento grupėse yra labai panašus ($M = 26,7439$, $M = 27,0303$, $M = 26,8765$, $M = 26,8571$), statistiškai reikšmingai nesiskiriantis. Tai pagrindžia eksperimento **grupių homogeniškumą amžiaus aspektu**.

Homogeniškumo patikrai kategorinis šalies kintamasis suskirstytas į stambesnes grupes (Lietuva, Jungtinė Karalystė ir kitos šalys), nes to reikalauja Chi kvadrato statistika, kuomet neturi būti tokių grupių, kuriose yra mažai stebėjimų. Gauti testo rezultatai atskleidžia, jog šalių proporcijos statistiškai reikšmingai nesiskiria ($p = 0,726$), tai patvirtina eksperimento **grupių homogeniškumą pagal dalyvių identifikaciją su šalimi**.

Skalių kintamiesiems – patirtis, naudojantis DI paremtais sprendimais ir naudojimosi plaukų priežiūros paslaugomis dažnis – taikytas vienos krypties tarpgrupinis ANOVA testas. Pastebima, jog vidurkiai, pagal naudojimąsi DI paremtais sprendimais patirtį, yra labai panašūs ($M = 2,88$, $M = 3,15$, $M = 3,22$, $M = 3,04$). Leveno testas rodo, kad dispersijų lygybė egzistuoja ($p = 0,672$). ANOVA testas ($p = 0,296$) leidžia pagrįsti, jog tarp grupių nėra statistiškai reikšmingų skirtumų, **pagal patirtį naudojantis DI paremtomis paslaugomis**. Tai patvirtina, jog pagal šį požymį **grupės yra homogeniškos**. Vidurkiai panašūs ir pagal naudojimosi grožio paslaugomis patirtį ($M = 2,7561$, $M = 2,8030$, $M = 2,8272$, $M = 2,7857$), tačiau Leveno testas atskleidžia, jog dispersijos nėra lygios ($p = 0,005$), todėl taikytas dispersijų nevienodumui atsparus Welch robustinis testas (Pallant, 2010). Testo rezultatai ($p=0,982$) leidžia pagrįsti, **pagal plaukų priežiūros paslaugų naudojimo dažnį, grupės yra homogeniškos** (žr. **13 priede**).

Nustatyta, kad polinkis priskirti moralumą reikšmingai skyrėsi tarp vyrų ir moterų. Ši išvada nurodo, jog moterys yra linkusios suvokti sprendimo priėmėjus, kaip labiau moralius nei vyrai. Apžvelgus skirtumus tarp respondentų grupių pagal išsilavinimą ir naudojimosi DI paremtomis paslaugomis patirtį, identifikuota, jog šios tiriamųjų charakteristikos neturėjo įtakos skirtumams tarp tiriamųjų kintamųjų vertinimo. Patvirtinta, jog visos eksperimento grupės pagal sociodemografinės charakteristikas yra homogeniškos.

4.4. Vartotojų polinkio priskirti moralinį veiksnumą, sprendimo palankumo, suvokiamo sąžiningumo ir pasitenkinimo sąsajų tyrimo rezultatų analizė

Koreliacinei analizei atlikti, svarbu įvertinti tiriamų kintamųjų normalumą (Pallant, 2010). Statistiniai vertinimo kriterijai, kaip Kolmogorovo-Smirnov (K-S) testas, leidžia įvertinti turimų duomenų sklaidą ir nustatyti, ar kintamieji statistiškai reikšmingai pasiskirstę apie normalųjį skirstinį. Duomenų normalumas tikrintas remiantis p reikšme (žr. 25 lentelė). Testo rezultatai pateikiami ir 14 priede.

25 lentelė. Kolmogorovo-Smirnov testo rezultatai (N=313)

Skalės pavadinimas	p reikšmė
Polinkio priskirti moralinį veiksnumą skalė (moralumas)	<0,001
Polinkio priskirti moralinį veiksnumą skalė (veiksnumas)	<0,001
Vartotojų suvokiamo sąžiningumo skalė	<0,001
Vartotojų pasitenkinimo skalė	<0,001

Gauta K-S statistika rodo, jog visi tiriami kintamieji yra statistiškai reikšmingai nutolę nuo normaliojo skirstinio. Šiuos rezultatus grindžia p reikšmės rezultatai, kuri visuose kintamuosiuose yra $< 0,001$, kai būtina normalumo sąlyga patenkinama $p > 0,05$ atveju (Pallant, 2010).

Remiantis normalumo patikros duomenimis, tolesnė ryšio stiprumo ir krypties analizė tarp kintamųjų, atliekama naudojant Spearmano koreliacijos koeficientą. Duomenys interpretuojami pagal Pallant (2010) pateiktus koreliacijos vertinimo rodiklius: kai $r =$ nuo 0,0 iki 0,09, koreliacija neegzistuoja, kai $r =$ nuo 0,10 iki 0,29 indikuojama maža koreliacija, kai $r =$ nuo 0,30 iki 0,49 fiksuojama vidutinio stiprumo koreliacija, kai $r =$ nuo 0,50 iki 1,00 egzistuoja stipri koreliacija (žr. 26 lentelė ir 15 priede).

26 lentelė. Vartotojų polinkio priskirti moralinį veiksnumą, sprendimo palankumo, suvokiamo sąžiningumo ir pasitenkinimo kintamųjų koreliacinė analizė (N=313)

	Polinkis priskirti moralinį veiksnumą (moralumas)	Polinkis priskirti moralinį veiksnumą (veiksnumas)	Vartotojų suvokiamas sąžiningumas	Vartotojų pasitenkinimas
Polinkis priskirti moralinį veiksnumą (moralumas)	1,000	-0,341**	0,631**	0,752**
Polinkis priskirti moralinį veiksnumą (veiksnumas)	-0,341**	1,000	-0,041	-0,218**
Vartotojų suvokiamas sąžiningumas	0,631**	-0,041	1,000	0,717**
Vartotojų pasitenkinimas	0,752**	-0,218**	0,717**	1,000

Pastaba: ** = $p < 0,001$

Spearmano koreliacijos analizės rezultatai atskleidžia, jog tarp daugelio tiriamų konstrukčių egzistuoja statistiškai reikšmingas ryšys. Stipriausias, statistiškai reikšmingas, teigiamas ryšys (0,752**) egzistuoja tarp polinkio priskirti moralumą ir vartotojo pasitenkinimo. Ryšio kryptis nurodo, jog didėjant suvokiamam moralumui, didėja vartotojų pasitenkinimas. Taip pat, stiprus ryšys (0,717**) nustatytas tarp vartotojų suvokiamo sąžiningumo ir pasitenkinimo. Teigiamas ryšio kryptis indikuoja, jog didėjant suvokiamam sąžiningumui, didėja vartotojų pasitenkinimas. Reikšmingai stipriai teigiamai koreliuoja (0,631**) vartotojų polinkis priskirti moralumą ir suvokiamas sąžiningumas. Suvokiamas moralumas teigiamai veikia suvokiamą sąžiningumą. Statistiškai reikšmingas, vidutinio stiprumo ryšys (-0,341**) nustatytas tarp vartotojų polinkio priskirti moralumą ir veiksnumą. Neigiama ryšio kryptis nurodo, jog didėjant polinkiui priskirti moralumą, mažėja polinkis priskirti veiksnumą ir atvirkščiai. Silpna, statistiškai reikšminga, neigiama koreliacija (-0,218**) stebima tarp vartotojų polinkio priskirti veiksnumą ir pasitenkinimo. Atvirkštinis ryšys suponuoja, jog didėjant suvokiamam veiksnumui, mažėja vartotojų pasitenkinimas ir atvirkščiai. Tarp polinkio priskirti veiksnumą ir suvokiamo sąžiningumo ryšys nenustatytas (-0,041).

Atsižvelgiant į vykdomo tyrimo metodiką, dviejų krypčių dispersinei analizei atlikti pasitelktas tarpgrupinis nepriklausomų imčių ANOVA (angl. *two-way between groups Anova*) metodas (žr. **16 priede**). Dvipusė ANOVA leidžia įvertinti nepriklausomų kintamųjų poveikį priklausomam kintamajam ir nustatyti galimą sąveikos efektą (Pallant, 2010). Šiame tyrime tikrinamas sprendimo priėmėjo (du lygmenys: DI ir žmogus) ir sprendimo palankumo (du lygmenys: palankus ir nepalankus) poveikis vartotojų polinkiui priskirti moralinį veiksnumą, pasitelkiant dvi moralinį veiksnumą reprezentuojančias dimensijas. Todėl, šis metodas pasitelktas hipotezėms: **H1a**, **H1b**, **H2a**, **H2b**, **H3a** ir **H3b** tikrinti:

H1a: *DI priimtas sprendimas (palyginus su žmogaus sprendimu) sukels polinkį priskirti didesnį moralumą sprendimo priėmėjui;*

H1b: *Žmogaus priimtas sprendimas (palyginus su DI sprendimu) sukels polinkį priskirti didesnį veiksnumą sprendimo priėmėjui;*

H2a: *Palankus sprendimas (palyginus su nepalankiu) sukels polinkį priskirti didesnį moralumą sprendimo priėmėjui;*

H2b: *Palankus sprendimas (palyginus su nepalankiu) sukels polinkį priskirti didesnį veiksnumą sprendimo priėmėjui;*

H3a: *Kai sprendimas yra palankus, nepriklausomai nuo to, ar sprendimą priėmė DI, ar žmogus, sprendimo priėmėjas bus suvokiamas, kaip labiau moralus;*

H3b: *Kai sprendimas yra nepalankus, DI sprendimų priėmėjas (palyginus su sprendimų priėmėju žmogumi) bus suvokiamas, kaip labiau veiksnus.*

Žemiau pateiktoje lentelėje, atskleidžiamas bendras tyrimo dalyvių pasiskirstymas pagal keturis tyrimo scenarijus: DI sprendimo priėmėjo, žmogaus sprendimo priėmėjo, palankaus sprendimo ir nepalankaus sprendimo (žr. **27 lentelė**).

27 lentelė. Tyrimo dalyvių pasiskirstymas pagal scenarijus (N=313)

Nepriklausomi kintamieji	Nepriklausomų kintamųjų lygmenys	N
Sprendimo priėmėjas	DI	163
	Žmogus	150
Sprendimo palankumas	Palankus	148
	Nepalankus	165

Siekiant užtikrinti dispersijos analizės patikimumą, atliktas Leveno variacijų lygybės testas, kuris nurodo ar dispersijos tarp kintamųjų yra vienodos (žr. **28 lentelė**).

28 lentelė. Lygių dispersijų prielaidos Leveno testas: sprendimo priėmėjo ir sprendimo palankumo poveikis vartotojų polinkiui priskirti moralinį veiksnumą (N=313)

Priklausomi kintamieji	Rodikliai	Levene statistika	df1	df2	p reikšmė
Polinkio priskirti moralinį veiksnumą (moralumas)	Remiantis vidurkiu	2,968	3	309	0,032
	Remiantis mediana	2,886	3	309	0,036
	Remiantis mediana ir pakoreguota df	2,886	3	308,095	0,036
	Remiantis nukirptu vidurkiu	3,074	3	309	0,028
Polinkio priskirti moralinį veiksnumą (veiksnumas)	Remiantis vidurkiu	,127	3	309	0,944
	Remiantis mediana	,169	3	309	0,917
	Remiantis mediana ir pakoreguota df	,169	3	294,098	0,917
	Remiantis nukirptu vidurkiu	,179	3	309	0,910

Gauti Leveno statistikos rezultatai rodo, jog dispersijų homogeniškumo prielaida nepažeista, todėl galima atlikti tolesnę analizę, norint patikrinti **H1a, H1b, H2a, H2b, H3a** ir **H3b** hipotezes (žr. **29 lentelė**).

29 lentelė. Sprendimo priėmėjo ir sprendimo palankumo poveikis priklausomiems kintamiesiems (N=313)

	Polinkio priskirti moralinį veiksnumą (moralumas) p reikšmė	Polinkio priskirti moralinį veiksnumą (veiksnumas) p reikšmė
Sprendimo priėmėjas	0,053	<0,001
Sprendimo palankumas	<0,001	<0,001
Sprendimo priėmėjas * sprendimo palankumas (moderavimo efektas)	0,606	0,196

Vertinant gautus rezultatus, galima teigti, jog sprendimo palankumas nemoderuoja sprendimo priėmėjo poveikio vartotojų polinkiui priskirti moralumą sprendimo priėmėjams – sprendimo

priėmėjo ir sprendimo palankumo sąveika nėra statistiškai reikšminga ($p = 0,606$). Toliau analizuojant atskirai nepriklausomų kintamųjų (sprendimo priėmėjo ir sprendimo palankumo) poveikį priskiriamam moralumui, matoma, kad sprendimo priėmėjo poveikis nėra statistiškai reikšmingas ($p = 0,053$). Rezultatai atskleidžia, kad sprendimo palankumas daro įtaką polinkiui priskirti moralumą ($p < 0,001$). Sprendimo priėmėjo ir sprendimo palankumo sąveika nėra statistiškai reikšminga – sprendimo palankumas nemoderuoja sprendimo priėmėjo poveikio polinkiui priskirti veiksnumą ($p = 0,196$). Toliau analizuojant atskirų nepriklausomų kintamųjų individualius efektus matoma, kad sprendimo priėmėjas daro statistiškai reikšmingą poveikį polinkiui priskirti veiksnumą, taip pat kaip ir sprendimo palankumas statistiškai reikšmingai veikia polinkį priskirti veiksnumą ($p < 0,001$). Remiantis šiais rezultatais, hipotezės **H1b**, **H2a** ir **H2b** patvirtinamos, o **H1a**, **H3a** ir **H3b** atmetamos.

Siekiant detalizuoti gautus rezultatus, pateikti dviejų krypčių dispersijos analizės rezultatai, kuria siekta nustatyti sprendimo priėmėjo ir sprendimo palankumo poveikį vartotojų polinkiui priskirti moralumą, tikrinant **H1a**, **H2a** ir **H3a** hipotezes (žr. **30 lentelė**).

30 lentelė. Tarpgrupinė nepriklausomų imčių dispersinė analizė: sprendimo priėmėjo ir sprendimo palankumo poveikio polinkiui priskirti moralumą testas (N=313)

Šaltinis	Kvadratų suma	Laisvės laipsniai (df)	Dispersijos įverčiai	Statistika (F)	p reikšmė	Dalinis Etakvadrato koeficientas
Koreguotas modelis	92,120 ^a	3	30,707	38,425	<0,001	0,272
Konstanta	1773,822	1	1773,822	2219,652	<0,001	0,878
Sprendimo priėmėjas	3,018	1	3,018	3,776	0,053	0,012
Sprendimo palankumas	89,895	1	89,895	112,489	<0,001	0,267
Sprendimo priėmėjo ir sprendimo palankumo sąveika	0,214	1	0,214	0,267	0,606	0,001

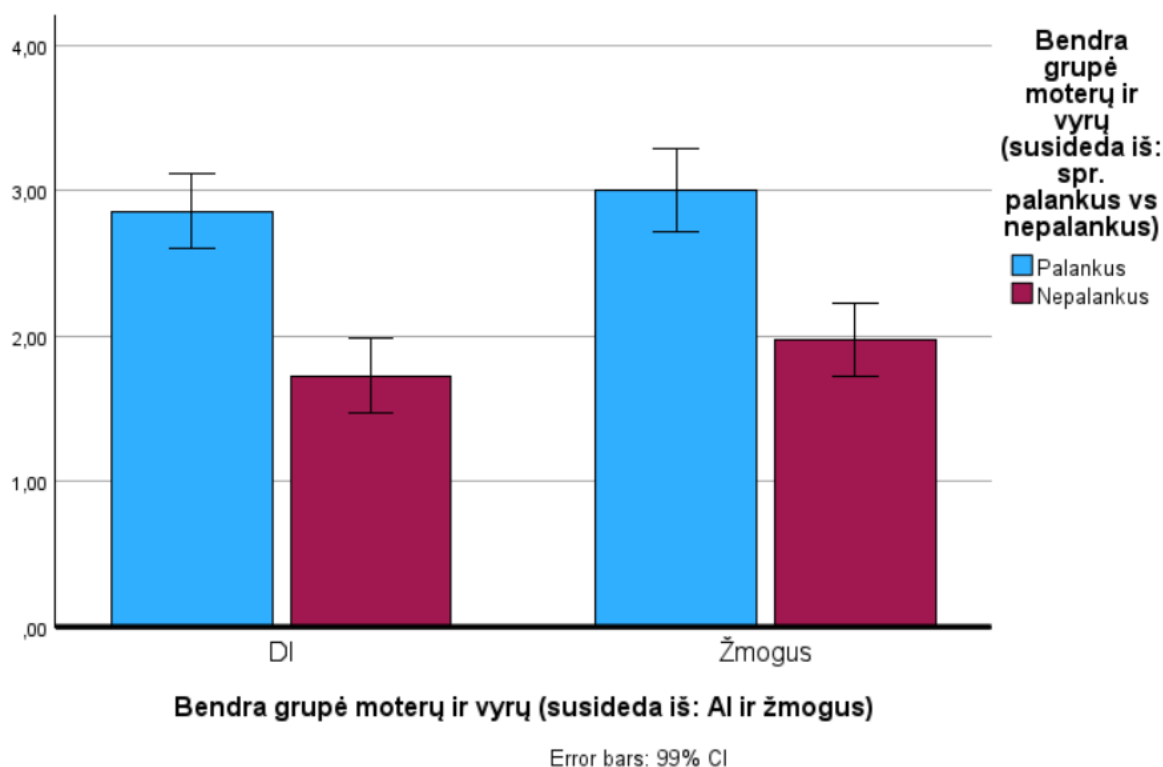
Rezultatų interpretacijos sumetimais, pateikiami ir polinkio priskirti moralumą, vidurkių palyginimai tarp skirtingų tyrimo scenarijų grupių (žr. **31 lentelė**).

31 lentelė. Polinkio priskirti moralumą vidurkių palyginimai skirtingų scenarijų grupėse (N=313)

Sprendimo priėmėjas	Sprendimo palankumas	Vidurkis	Standartinis nuokrypis (SN)	N
DI	Palankus	2,8577	0,96407	82
	Nepalankus	1,7284	0,80728	81
	Bendras	2,2965	1,05231	163
Žmogus	Palankus	3,0025	0,78146	66
	Nepalankus	1,9782	0,98066	84
	Bendras	2,4289	1,03073	150

Bendras	Palankus	2,9223	0,88735	148
	Nepalankus	1,8556	0,90571	165
	Bendras	2,3600	1,04246	313

Žemiau pateikiama histograma, iliustruojanti individualius sprendimo priėmėjo ir sprendimo palankumo poveikius polinkiui priskirti moralumą bei šių kintamųjų sąveikos poveikis (žr. **3 pav.**).



3 pav. Polinkio priskirti moralumą vidurkių palyginimas tarp nepriklausomų grupių

Gauti duomenys atskleidžia, jog sprendimo palankumo ir sprendimo priėmėjo sąveika polinkiui priskirti moralumą nėra statistiškai reikšminga, $F(1, 309) = 0,267$, $p = 0,606$ (žr. **30 lentelė**).

Statistiškai reikšmingų skirtumų polinkio priskirti moralumą vidurkiuose, skirtingose sprendimo priėmėjo grupėse, nebuvo nustatyta ($F(1, 309) = 3,018$, $p = 0,053$) (žr. **30 lentelė**). Stebimi panašūs polinkio priskirti moralumą vidurkiai, DI sprendimo priėmėjo ($M = 2,2965$, $SN = 1,05231$) ir žmogaus sprendimo priėmėjo ($M = 2,4289$, $SN = 1,03073$) grupėse (žr. **31 lentelė**). Šie rezultatai rodo, jog polinkis priskirti moralumą nėra veikiamas sprendimo priėmėjo tipo.

Tačiau, stebimas statistiškai reikšmingas sprendimo palankumo poveikis polinkiui priskirti moralumą: $F(1, 309) = 112,489$, $p < 0,001$ (žr. **30 lentelė**). Kai sprendimas palankus, polinkio priskirti moralumą vidurkis ($M = 2,9223$, $SN = 0,88735$) yra reikšmingai didesnis, lyginant su nepalankaus sprendimo grupės, polinkio priskirti moralumą vidurkiu ($M = 1,8556$, $SN = 0,90571$) (žr. **31 lentelė**). Kuomet sprendimą priima DI, polinkio priskirti moralumą vidurkis yra didesnis palankaus sprendimo atveju ($M = 2,8577$, $SN = 0,96407$) nei nepalankaus sprendimo ($M = 1,7284$, $SN = 0,80728$). Panašios tendencijos stebimos ir, kai sprendimą priima žmogus – polinkio priskirti moralumą vidurkis yra aukštesnis esant palankiam sprendimui ($M = 0,80728$, $SN = 0,78146$), lyginant su nepalankaus sprendimo grupe ($M = 1,9782$, $SN = 0,98066$). Sprendimo palankumo poveikis

polinkiui priskirti moralumą yra stiprus, nes dalinio Eta-kvadrato koeficientas yra didesnis nei 0,138 (Pallant, 2010) ir lygus $2\eta = 0,267$.

Pagal gautus duomenis galima matyti, jog sprendimo palankumas teigiamai veikia polinkį priskirti moralumą, nepriklausomai nuo sprendimo priėmėjo tipo. Polinkis priskirti moralumą išliko panašus nepriklausomai nuo to, ar sprendimą priėmė DI, ar žmogus, o sprendimo priėmėjo poveikio polinkiui priskirti moralumą stiprumas (stebima nuosekli polinkio priskirti moralumą dinamika skirtingose sprendimo priėmėjo grupėse) reikšmingai nekinta. Šią išvadą grindžia ir neegzistuojantis sprendimo priėmėjo ir sprendimo palankumo sąveikos efektas (žr. **30 lentelė**). Sprendimo palankumas neturi poveikio tam, kaip sprendimo priėmėjas veikia polinkį priskirti moralumą. Lygiai taip pat sprendimo priėmėjo tipas neturi poveikio tam, kaip sprendimo palankumas veikia polinkį priskirti moralumą. Polinkis priskirti moralumą nepriklauso nuo to, ar sprendimo priėmėjas yra žmogus ar DI. Tuo tarpu sprendimo palankumas veikia polinkį priskirti moralumą sprendimo priėmėjui. Palankaus sprendimo atveju, vartotojai sprendimo priėmėjui priskiria didesnę moralumą nei nepalankaus sprendimo atveju.

Atsižvelgiant į analizės rezultatus, hipotezės **H1a**: *DI priimtas sprendimas (palyginus su žmogaus sprendimu) sukels polinkį priskirti didesnę moralumą sprendimo priėmėjui*; ir **H3a**: *Kai sprendimas yra palankus, nepriklausomai nuo to, ar sprendimą priėmė DI, ar žmogus, sprendimo priėmėjas bus suvokiamas, kaip labiau moralus*; **nepatvirtinamos**. Hipotezė **H2a**: *Palankus sprendimas (palyginus su nepalankiu) sukels polinkį priskirti didesnę moralumą sprendimo priėmėjui* **patvirtinama**.

Toliau, tarpgrupinė dviejų krypčių dispersija atlikta, siekiant įvertinti sprendimo priėmėjo ir sprendimo palankumo poveikį vartotojų polinkiui priskirti veiksnumą, tikrinant **H1b**, **H2b** ir **H3b** hipotezes (žr. **32 lentelė**).

32 lentelė. Tarpgrupinė nepriklausomų imčių dispersinė analizė: sprendimo priėmėjo ir sprendimo palankumo poveikio polinkiui priskirti veiksnumą testas (N=313)

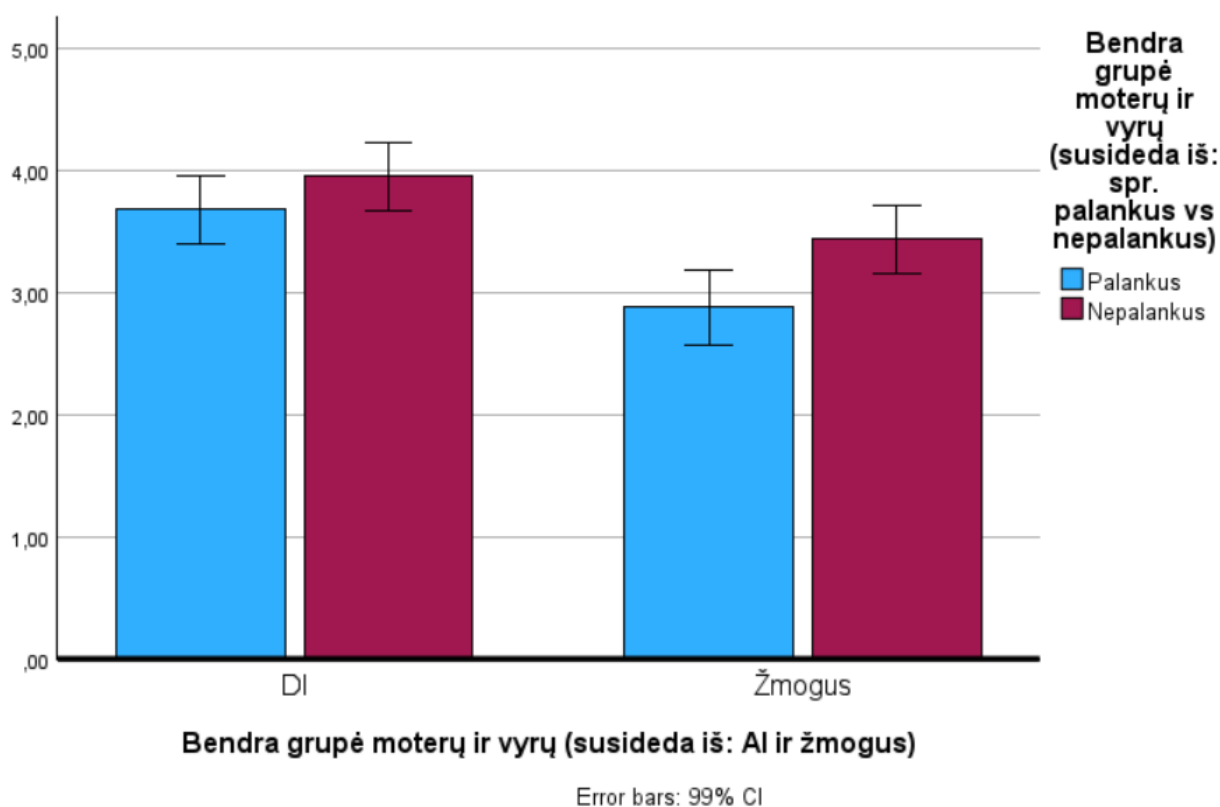
Šaltinis	Kvadratų suma	Laisvės laipsniai (df)	Dispersijos įverčiai	Statistika (F)	p reikšmė	Dalinis Etakvadrato koeficientas
Koreguotas modelis	44,858 ^a	3	14,953	15,960	<0,001	0,134
Konstanta	3768,232	1	3768,232	4021,954	<0,001	0,929
Sprendimo priėmėjas	33,625	1	33,625	35,889	<0,001	0,104
Sprendimo palankumas	13,238	1	13,238	14,130	<0,001	0,044
Sprendimo priėmėjo ir sprendimo palankumo sąveika	1,574	1	0,937	1,680	0,196	0,005

Norint geriau interpretuoti gautus duomenis, **33 lentelė** pristato vidurkių palyginimus tarp skirtingų tyrimo scenarijų.

33 lentelė. Polinkio priskirti veiksnumą vidurkių palyginimai skirtingų scenarijų grupėse (N=313)

Sprendimo priėmėjas	Sprendimo palankumas	Vidurkis	Standartinis nuokrypis (SN)	N
DI	Palankus	3,6799	0,93068	82
	Nepalankus	3,9506	1,01506	81
	Bendras	3,8144	0,97996	163
Žmogus	Palankus	2,8788	0,96201	66
	Nepalankus	3,4345	0,96171	84
	Bendras	3,1900	0,99777	150
Bendras	Palankus	3,3226	1,02281	148
	Nepalankus	3,6879	1,01865	165
	Bendras	3,5152	1,03522	313

Vidurkių palyginimo rezultatai grafiškai atvaizduojami **4 pav.**



4 pav. Polinkio priskirti veiksnumą vidurkių palyginimas tarp nepriklausomų grupių

Remiantis $F(1, 309) = 1,680$ ir $p = 0,196$ reikšmėmis, patvirtinama, jog sprendimo palankumo ir sprendimo priėmėjo sąveikos efekto dydžio poveikis, polinkiui priskirti veiksnumą, nepasižymi statistiniu reikšmingumu (žr. **32 lentelė**)

Pastebima statistiškai reikšminga polinkio priskirti veiksnumą priklausomybė nuo sprendimo priėmėjo ($F(1, 309) = 35,889$, $p < 0,001$) (žr. **32 lentelė**). Polinkiui priskirti veiksnumą, palankaus sprendimo situacijoje, nustatyta aukštesnė vidurkio reikšmė DI sprendimo priėmėjo atžvilgiu ($M =$

3,6799, SN = ,93068) nei sprendimo priėmėjo žmogaus (M = 2,8788, SN = 0,96201) (žr. **33 lentelė**). Polinkis priskirti veiksnumą reikšmingai skiriasi ir nepalankaus sprendimo situacijoje, DI sprendimo priėmėjo vidurkis yra aukštesnis (M = 3,9506, SN = 1,01506), lyginant su žmogaus sprendimo priėmėjo vidurkiu (M = 3,4345, SN = 0,96171). Pagal naudojamą skalę, aukštesnis vidurkis reiškia didesnę priklausomybę nuo kitų, o tai reprezentuoja mažesnę veiksnumą. Šie rezultatai rodo, jog tyrimo dalyviai yra linkę suvokti DI, kaip labiau priklausomą ir mažiau veiksnų nei žmogų.

Įvertinus gautus rezultatus, pastebima, jog sprendimo priėmėjas turi teigiamą poveikį polinkiui priskirti veiksnumą, nepriklausomai nuo sprendimo palankumo. Nors pastebėta, jog polinkis priskirti mažesnę veiksnumą yra aukštesnis nepalankaus sprendimo atveju, kai sprendimą priima DI, lyginant su žmogumi, sprendimo priėmėjo poveikio polinkiui priskirti veiksnumą stiprumas (matoma tolygi polinkio priskirti veiksnumą dinamika skirtingose sprendimo palankumo grupėse) reikšmingai nesiskiria. Visa tai paaiškina ir statistiškai nereikšmingas sprendimo priėmėjo ir sprendimo palankumo sąveikos efektas (žr. **32 lentelė**). Sprendimo palankumas neveikia to, kaip sprendimo priėmėjas veikia polinkį priskirti veiksnumą. Pastebima, jog sprendimo priėmėjo efektas polinkiui priskirti veiksnumą yra didelis, $2\eta = ,104$.

Be to, polinkis priskirti veiksnumą reikšmingai skiriasi sprendimo palankumo grupėse (F (1, 309) = 35,889, $p < 0,001$) (žr. **32 lentelė**). Nepalankaus sprendimo grupėje pastebėtas didesnis polinkis priskirti mažesnę veiksnumą (M = 3,6879, SN = 1,02281), lyginant su palankaus sprendimo situacija (M = 3,3226, SN = 1,01865) (žr. **33 lentelė**). Kai sprendimą priima DI, mažesnis veiksnumas (didesnė priklausomybė) priskiriamas nepalankaus sprendimo atveju (M = 3,9506, SN = 1,01506), nei palankaus (M = 3,6799, SN = 0,93068). Stebimi panašūs rezultatai ir kai sprendimą priima žmogus, mažesnis polinkis priskirti veiksnumą (didesnė priklausomybė) išryškėja nepalankaus sprendimo situacijoje (M = 3,4345, SN = 0,96171) nei palankaus (M = 2,8788, SN = 0,96201). Dalinio eta kvadrato reikšmė indikuoja, jog nepaisant to, kad sprendimo palankumas teigiamai veikia polinkį priskirti veiksnumą, efektas yra silpnas ($2\eta = 0,044$).

Remiantis gautais rezultatais, hipotezės **H1b**: *Žmogaus priimtas sprendimas (palyginus su DI sprendimu) sukelia polinkį priskirti didesnę veiksnumą sprendimo priėmėjui*; ir **H2b**: *Palankus sprendimas (palyginus su nepalankiu) sukelia polinkį priskirti didesnę veiksnumą sprendimo priėmėjui patvirtinamos*. Hipotezė **H3b**: *Kai sprendimas yra nepalankus, DI sprendimų priėmėjas (palyginus su sprendimų priėmėju žmogumi) bus suvokiamas, kaip labiau veiksnus, nepatvirtinama*.

Siekiant patikrinti **H4a**, **H4b**, **H5a**, **H5b**, **H6**, **H6a**, **H7a** ir **H7b** hipotezes, atliekamos tiesinių regresijų analizės (žr. **17 priede**). Norint patikrinti ryšį tarp polinkio priskirti moralinį veiksnumą ir suvokiamo sąžiningumo, regresine analize tikrinamos **H4a** ir **H4b** hipotezės (žr. **34 lentelė** ir **35 lentelė**).

34 lentelė. Daugialypės tiesinės regresijos, tiriant ryšį tarp polinkio priskirti moralinį veiksnumą ir suvokiamo sąžiningumo, priklausomo kintamojo rezultatai (N=313)

Priklausomas kintamasis	Koreliacijos koeficientas (R)	Determinacijos koeficientas (R ²)	ANOVA	
			Statistika (F)	p reikšmė
Suvokiamas sąžiningumas	,628 ^a	,390	100,836	<,001 ^b

Žemiau pateikiami nepriklausomų kintamųjų analizės rezultatai (žr. **35 lentelė**).

35 lentelė. Daugialypės tiesinės regresijos, tiriant ryšį tarp polinkio priskirti moralinį veiksnumą ir suvokiamo sąžiningumo, nepriklausomų kintamųjų rezultatai

Nepriklausomas kintamasis	Nestand. koef.		Stand. koef.	p reikšmė	Kolinearumo statistika	
	B	Stand. paklaida	Beta (β)		Tolerancija	VIF
Polinkis priskirti moralinį veiksnumą (moralumas)	0,721	0,051	0,657	<0,001	0,912	1,096
Polinkis priskirti moralinį veiksnumą (veiksnumas)	0,234	0,051	0,212	<0,001	0,912	1,096

Kolinearumo statistika atskleidžia, jog modelyje nėra multikolinearumo problemos, nes tolerancijos reikšmė yra didesnė nei 0,10, o variacijos infliacijos koeficientas mažesnis nei 10 (Pallant, 2010): abejuose kintamuosiuose tolerancija = 0,912, VIF = 1,096. Gauti regresinės analizės rezultatai rodo, kad modelis yra statistiškai reikšmingas, $F = 100,836$, $p < 0,001$ (žr. **34 lentelė**). Nepriklausomi kintamieji moralumas ($p < 0,001$) ir veiksnumas ($p < 0,001$) daro statistiškai reikšmingą poveikį priklausomam, suvokiamo sąžiningumo kintamajam (žr. **35 lentelė**). Determinacijos koeficiento reikšmė ($R^2 = 0,390$) indikuoja, jog modelis paaiškina 39% suvokiamo sąžiningumo dispersijos. Galima matyti, jog polinkis priskirti moralumą, stipriai teigiamai veikia sąžiningumą: $B = 0,721$, $\beta = 0,657$, $p < 0,001$. Ši išvada rodo, jog kuo aukštesnis polinkis priskirti moralumą, tuo didesnis suvokiamas sąžiningumas. Panašus teigiamas, tačiau silpnesnis poveikis stebimas ir polinkio priskirti veiksnumą suvokiamo sąžiningumo atveju: $B = 0,234$, $\beta = 0,212$, $p < 0,001$.

Galima teigti, jog abu nepriklausomi kintamieji statistiškai reikšmingai teigiamai veikia suvokiamą sąžiningumą, tačiau polinkis priskirti moralumą, lyginant su polinkiu priskirti veiksnumą, stipriau prognozuoja gautą rezultatą. Remiantis rezultatais, polinkis priskirti moralinį veiksnumą teigiamai veikia suvokiamą sąžiningumą, todėl hipotezės **H4a**: *Vartotojų polinkis priskirti moralumą sprendimo priėmėjui daro teigiamą poveikį suvokiamam sąžiningumui*, ir **H4b**: *Vartotojų polinkis priskirti veiksnumą sprendimo priėmėjui daro teigiamą poveikį suvokiamam sąžiningumui, patvirtinamos*.

Siekiant nustatyti, kaip šį ryšį moderuoja sprendimo priėmėjas, hipotezėms **H5a** ir **H5b** patikrinti atlikta regresinė analizė su moderavimo efektu, naudojant Hayes'o *PROCESS* v5.0 makrokomandos, skirtos *SPSS IBM* programinei įrangai, modelį Nr. 1.

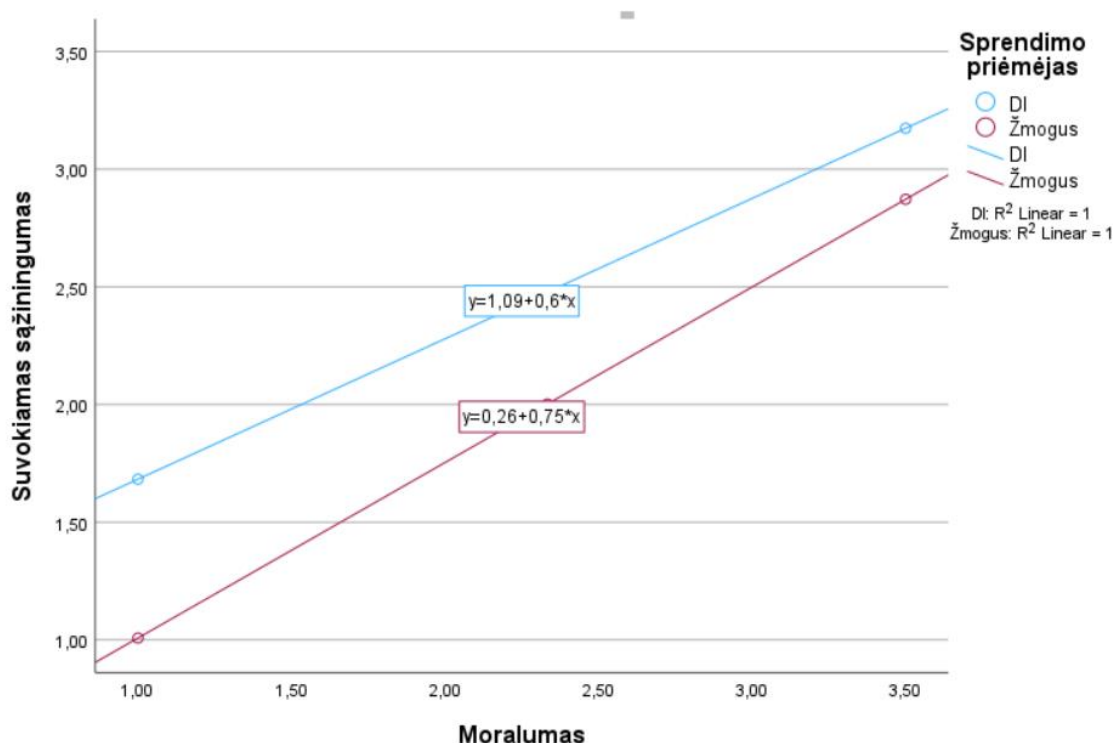
Pristatomi sprendimo priėmėjo ir moralumo sąveikos poveikio rezultatai suvokiamam sąžiningumui, norint patikrinti hipotezę **H5a** (žr. **36 lentelė**).

36 lentelė. Regresijos modelių rezultatai, tiriant polinkio priskirti moralumą ir sprendimo priėmėjo sąveikos poveikį suvokiamam sąžiningumui (N=313)

R	Priklausomas kintamasis			
	Y: Suvokiamas sąžiningumas			
	→	Koef.	SE	p

Konstanta	$b_0 \rightarrow$	1,9103	0,3841	0,00000
X: Moralumas	$b_1 \rightarrow$	0,4473	0,1506	0,0032
W: Sprendimo priėmėjas	$b_2 \rightarrow$	-0,8242	0,2510	0,0011
X*W	$b_3 \rightarrow$	0,1493	0,0971	0,1252
Modelis	$R^2 = 0,4000$; $F = 68,6760$; $p = 0,0000$			

Gauti rezultatai atskleidžia, jog prognostinis modelis pasižymi statistiniu reikšmingumu ($p = 0,0000$) ir paaiškina 40% ($R^2 = 0,4000$) suvokiamo sąžiningumo sklaidos apie vidurkį. Vis dėlto, statistiškai reikšmingas sprendimo priėmėjo poveikis ryšiui tarp polinkio priskirti moralumą ir suvokiamo sąžiningumo nenustatytas ($p = 0,1252$). Gauti rezultatai vizualizuoti **5 pav.**



5 pav. Polinkio priskirti moralumą poveikis suvokiamam sąžiningumui, priklausomai nuo sprendimo priėmėjo tipo

Remiantis analizės rezultatais, hipotezė **H5a**: *Sprendimo priėmėjas moderuoja ryšį tarp polinkio priskirti moralumą ir sąžiningumo: kai sprendimą priima žmogus, ryšys tarp moralumo ir sąžiningumo bus stipresnis palyginus su situacija, kai sprendimą priima DI, nepatvirtinama.*

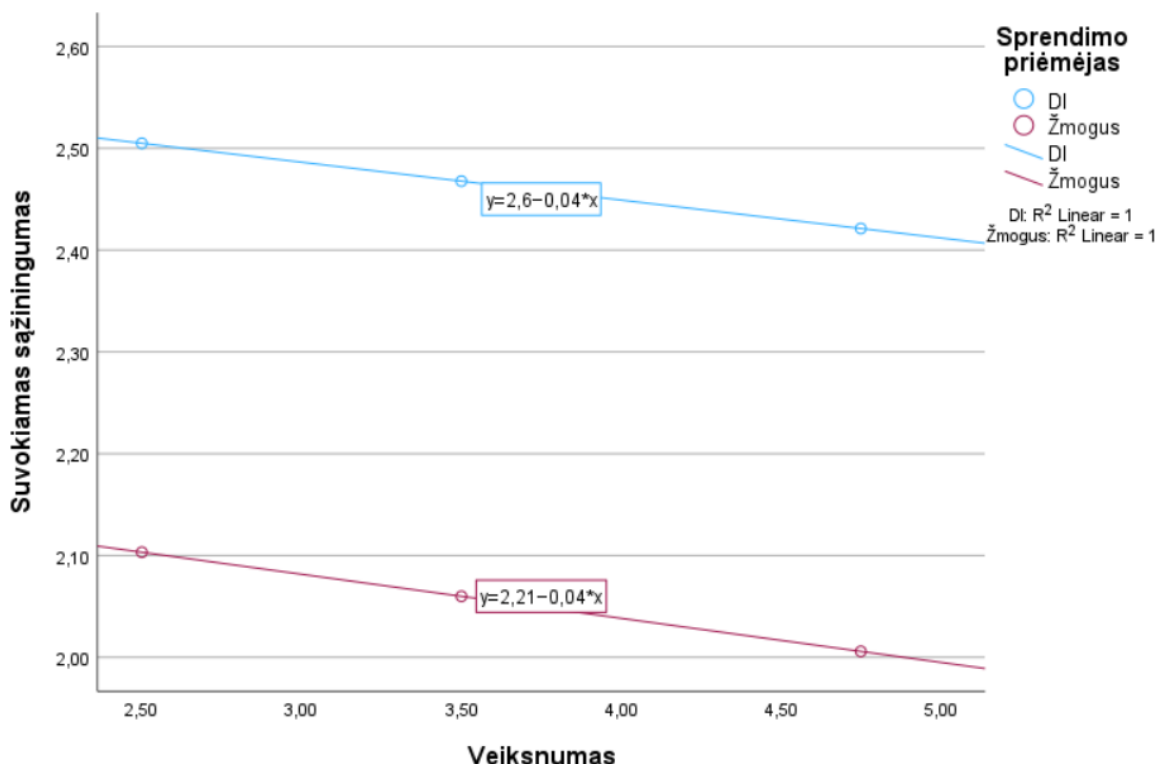
Toliau, pateikiami sprendimo priėmėjo ir polinkio priskirti veiksnų sąveikos poveikio suvokiamam sąžiningumui rezultatai, siekiant patikrinti **H5b** (žr. **37 lentelė**).

37 lentelė. Regresijos modelių rezultatai, tiriant polinkio priskirti veiksnų ir sprendimo priėmėjo sąveikos poveikį suvokiamam sąžiningumui (N=313)

R	Priklausomas kintamasis			
	Y: Suvokiamas sąžiningumas			
	$\rightarrow$	Koef.	SE	p
Konstanta	$b_0 \rightarrow$	2,9844	0,7797	0,0002

X: Veiksnumas	$b_1 \rightarrow$	-0,0311	0,2040	0,8789
W: Sprendimo priėmėjas	$b_2 \rightarrow$	-0,3865	0,4737	0,4152
$X*W$	$b_3 \rightarrow$	-0,0061	0,1300	0,9626
Modelis	$R^2 = 0,0292$; $F = 3,1013$; $p = 0,0270$			

Analizės rezultatai atskleidžia, jog sprendimo priėmėjas nemoderuoja polinkio priskirti veiksnų ryšio su suvokiamu sąžiningumu. Šios išvados grindžiamos prognostinio modelio statistiniu nereikšmingumu ($p = 0,0270$; $R^2 = 0,0292$) ir polinkio priskirti veiksnų, ir sprendimo priėmėjo sąveikos efekto nereikšmingumu ($p = 0,9626$). Gauti duomenys iliustruoti **6 pav.**



6 pav. Polinkio priskirti veiksnų poveikis suvokiamam sąžiningumui, priklausomai nuo sprendimo priėmėjo tipo

Atsižvelgiant į gautus rezultatus, hipotezė **H5b**: *Sprendimo priėmėjas moderuoja ryšį tarp polinkio priskirti veiksnų ir sąžiningumo: kai sprendimą priima DI, ryšys tarp veiksnų ir sąžiningumo bus stipresnis palyginus su situacija, kai sprendimą priima žmogus, nepatvirtinama.*

Toliau, atlikta tiesinė regresija, siekiant įvertinti galimą suvokiamo sąžiningumo ir pasitenkinimo ryšį, ir patikrinti **H6** hipotezę (žr. **38 lentelė**).

38 lentelė. Tiesinės regresijos, tiriant ryšį tarp suvokiamo sąžiningumo ir pasitenkinimo, kaip priklausomo kintamojo rezultatai (N=313)

Priklausomas kintamasis	Koreliacijos koeficientas (R)	Determinacijos koeficientas (R^2)	ANOVA	
			Statistika (F)	p reikšmė
Pasitenkinimas	0,669 ^a	0,448	252,076	<0,001 ^b

Žemiau pateikiami nepriklausomojo kintamojo rezultatai (žr. **39 lentelė**).

39 lentelė. Daugialypės tiesinės regresijos, tiriant ryšį tarp suvokiamo sąžiningumo ir pasitenkinimo, nepriklausomo kintamojo rezultatai

Nepriklausomas kintamasis	Nestand. koef.		Stand. koef.	p reikšmė
	B	Stand. paklaida	Beta (β)	
Suvokiamas sąžiningumas	0,782	0,049	0,669	<0,001

Regresijos modelis yra statistiškai reikšmingas ($F = 252,076$, $p < 0,001$.) Suvokiamas sąžiningumas reikšmingai veikia pasitenkinimą, $p < ,001$. Determinacijos koeficientas ($R^2 = 0,669$) atskleidžia, jog modelis paaiškina 44,8% pasitenkinimo dispersijos. Stebimas didelis suvokiamo sąžiningumo poveikis pasitenkinimui ($B = 0,782$, $\beta = 0,669$, $p < 0,001$). Gauti rezultatai nurodo, jog tarp kintamųjų egzistuoja stipri tiesioginė priklausomybė, didėjant suvokiamam sąžiningumui, didėja pasitenkinimas.

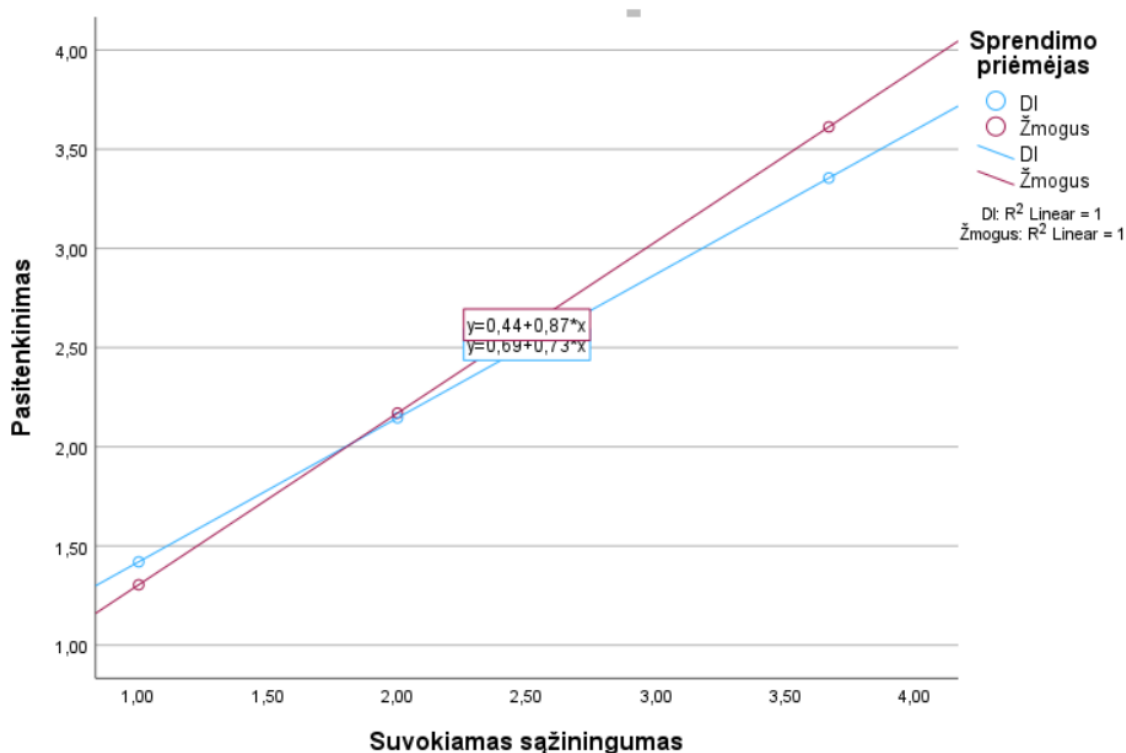
Rezultatai leidžia teigti, jog suvokiamas sąžiningumas turi teigiamą poveikį pasitenkinimui, todėl hipotezė **H6**: *Suvokiamas sprendimo priėmėjo sąžiningumas teigiamai veikia vartotojų pasitenkinimą paslauga, patvirtinama.*

Norint įvertinti ar sprendimo priėmėjo tipas – DI arba žmogus, sąlygoja šio poveikio pokyčius, atlikta regresinė analizė, naudojant Hayes'o *PROCESS* makrokomandos moderavimo efekto modelį (Nr. 1). Toliau pateikiami sprendimo priėmėjo ir suvokiamo sąžiningumo sąveikos efektas pasitenkinimui, tikrinant **H6a** hipotezę (žr. **40 lentelė** ir **17 priede**)

40 lentelė. Regresijos modelių rezultatai, tiriant sprendimo priėmėjo ir suvokiamo sąžiningumo poveikį pasitenkinimui (N=313)

R	Priklausomas kintamasis			
	Y: Pasitenkinimas			
	→	Koef.	SE	p
Konstanta	$b_0 \rightarrow$	0,9500	0,4027	0,0189
X: Suvokiamas sąžiningumas	$b_1 \rightarrow$	0,5856	0,1530	0,0002
W: Sprendimo priėmėjas	$b_2 \rightarrow$	-0,2557	0,2530	0,3130
X*W	$b_3 \rightarrow$	0,1400	0,1008	0,1659
Modelis	$R^2 = 0,4516$; $F = 48,8041$; $p = 0,0000$			

Remiantis gautais rezultatais, sprendimo priėmėjas neturi statistiškai reikšmingo poveikio suvokiamo sąžiningumo ir pasitenkinimo ryšiui, $p = 0,1659$. Gauti rezultatai pavaizduoti **7 pav.**



7 pav. Suvokiamo sąžiningumo poveikis pasitenkinimui paslauga, priklausomai nuo sprendimo priėmėjo tipo

Todėl, hipotezė **H6a**: *Sprendimo priėmėjas moderuoja ryšį tarp sąžiningumo ir pasitenkinimo taip, kad: kai sprendimą priima DI, ryšys tarp sąžiningumo ir pasitenkinimo bus stipresnis, palyginus su situacija, kai sprendimą priima žmogus, atmetama.*

Vertinant ar suvokiamas sąžiningumas tarpininkauja ryšiui tarp polinkio priskirti moralinį veiksnumą ir vartotojų pasitenkinimo, atlikta keletą regresijų analizė, taikant Hayes'o *PROCESS* makrokomandos mediacijos efekto modelį Nr. 4 (žr. **17 priede**). Pirmiausia, siekta nustatyti, kaip suvokiamas sąžiningumas paaiškina moralumo poveikį pasitenkinimui, tikrinant **H7a** hipotezę (žr. **41 lentelė** ir **42 lentelė**).

41 lentelė. Regresijos modelių rezultatai, tiriant ryšį tarp suvokiamo sąžiningumo, moralumo ir pasitenkinimo (N=313)

R	Priklausomas kintamasis											
	M: S				Y:P				Y:P			
	→	Koef.	SE	p	→	Koef.	SE	p	→	Koef.	SE	p
Konstanta	i_M →	0,7336	0,1291	0,0000	i_Y →	-0,1949	0,1195	0,1039	i_Y →	0,1056	0,1253	0,4000
X: M	a →	0,6522	0,0500	0,0000	c' →	0,6881	0,0548	0,0000	c →	0,9553	0,0486	0,0000
M: S	-	-	-	-	b →	0,4096	0,0500	0,0000	-	-	-	-
Modelis	$R^2 = 0,3532$; $F = 169,8179$; $p = \mathbf{0,0000}$				$R^2 = ,7961$; $F = 268,2145$; $p = \mathbf{0,0000}$				$R^2 = ,7446$; $F = 386,8905$; $p = \mathbf{0,0000}$			

Pastaba: M – moralumas (X), S – suvokiamas sąžiningumas (M), P – pasitenkinimas (Y).

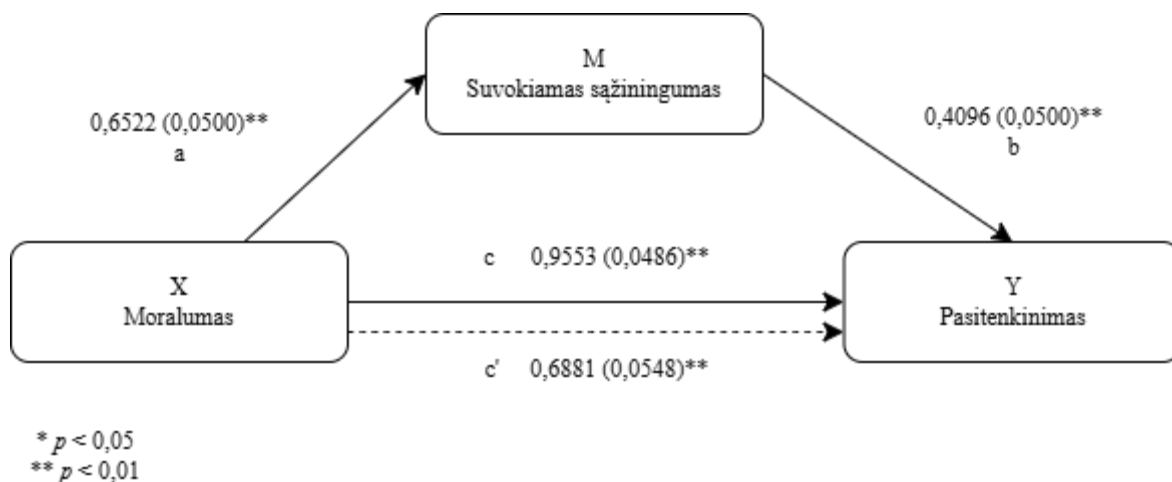
Duomenys rodo, jog moralumo poveikio suvokiamam sąžiningumui modelis yra statistiškai reikšmingas ($R^2 = 0,3532$, $p = 0,0000$). Interpretuojant rezultatus, patvirtinama, kad moralumas reikšmingai teigiamai prognozavo suvokiamą sąžiningumą (a kelias = 0,6522, $p = 0,0000$). Toliau, įvertinus moralumo ir suvokiamo sąžiningumo įtaką pasitenkinimui modelio statistinį reikšmingumą ($R^2 = 0,7961$, $p = 0,0000$), galima teigti, jog abu kintamieji reikšmingai teigiamai veikia pasitenkinimą (c' kelias = 0,6881, $p = 0,0000$; b kelias = 0,4096, $p = 0,0000$). Analizė patvirtino moralumo poveikio pasitenkinimui modelio statistinį reikšmingumą ($R^2 = 0,7446$, $p = 0,0000$) ir teigiamą reikšmingą įtaką (c kelias = 0,9553, $p = 0,0000$).

Žemiau pateikiami tiesioginės, netiesioginės ir suminės moralumo įtakos pasitenkinimui rezultatai (žr. **42 lentelė**). Poveikis patvirtinamas tuomet, kai pasikliautiniojo intervalo reikšmė neapima nulinės reikšmės (Sürücü et al., 2023)

42 lentelė. Tiesioginė, netiesioginė (per suvokiamą sąžiningumą) ir suminė moralumo įtaka pasitenkinimui (N=313)

Kelias	Efektas (EF)	95% pasikliautinis intervalas*	
		LLCI	ULCI
Tiesioginė polinkio priskirti moralumą įtaka pasitenkinimui			
M → P (c‘)	0,6881	0,5803	0,7960
Netiesioginė polinkio priskirti moralumą įtaka pasitenkinimui			
M → S → P (a*b)	0,2671	0,1684	0,3842
Suminė polinkio priskirti moralumą įtaka pasitenkinimui			
M → P (c)	0,9553	0,8597	1,0508

Remiantis gautais rezultatais, stebimas stipresnis tiesioginis moralumo poveikis pasitenkinimui (EF = 0,6881, LLCI = 0,5803, ULCI = 0,7960), lyginant su netiesioginiu poveikiu per sąžiningumą (EF = 0,2671, LLCI = 0,1684, ULCI = 0,3842). Nepaisant to, pasikliautiniojo intervalo apatinė ir viršutinė ribos rodo, kad suvokiamas sąžiningumas turi mediavimo poveikį. Rezultatai vizualizuoti **8 pav.**



8 pav. Moralumo įtaka pasitenkinimui paslauga per suvokiamą sąžiningumą

Remiantis analizės rezultatais, hipotezė **H7a**: *Suvokiamas sąžiningumas medijuoja ryšį tarp moralumo ir pasitenkinimo taip, kad: moralumas daro teigiamą poveikį sąžiningumui, kuris savo ruožtu teigiamai veikia pasitenkinimą, patvirtinama.*

Toliau tikrinama hipotezė **H7b**, taikant tuos pačius, kelių regresijos analizės metodus (žr. **43 lentelė** ir **44 lentelė**).

43 lentelė. Regresijos modelių rezultatai, tiriant ryšį tarp suvokiamo sąžiningumo, veiksnų ir pasitenkinimo (N=313)

R	Priklausomas kintamasis											
	M: S				Y:P				Y:P			
	→	Koef.	SE	p	→	Koef.	SE	p	→	Koef.	SE	p
Konstanta	$i_M \rightarrow$	2,2048	0,1291	0,0000	$i_Y \rightarrow$	1,4939	0,2189	0,0000	$i_Y \rightarrow$	3,2277	0,2635	0,0000
X: V	$a \rightarrow$	0,0193	0,0627	0,7582	$c' \rightarrow$	-0,2620	0,0525	0,0000	$c \rightarrow$	-0,2469	0,0719	0,0007
M: S	-	-	-	-	$b \rightarrow$	0,7864	0,0475	0,0000	-	-	-	-
Modelis	$R^2 = 0,0175$; $F = 0,0949$; $p = 0,7582$				$R^2 = 0,6991$; $F = 148,2104$; $p = \mathbf{0,0000}$				$R^2 = 0,1911$; $F = 11,7851$; $p = \mathbf{0,0007}$			

Pastaba: V – veiksnas, S – sąžiningumas, P – pasitenkinimas.

Pastebima, kad veiksnų poveikio sąžiningumo modelis nėra statistiškai reikšmingas ($R^2 = 0,0175$, $p = 0,7582$), todėl šio ryšio duomenys toliau neinterpretuojami. Modelis, paaiškinantis veiksnų ir sąžiningumo poveikį statistiškai reikšmingas ($R^2 = 0,6991$; $p = 0,0000$). Veiksnas (c' kėias = -0,2620, $p = 0,0000$) ir sąžiningumas (b kėias = 0,7864, $p = 0,0000$) abu reikšmingai veikia pasitenkinimą. Veiksnų poveikio pasitenkinimui modelis yra reikšmingas ($R^2 = 0,1911$; $p = 0,0007$), atskleidžiantis, jog veiksnas reikšmingai teigiamai veikia pasitenkinimą (c kėias = -0,2469, $p = 0,0007$).

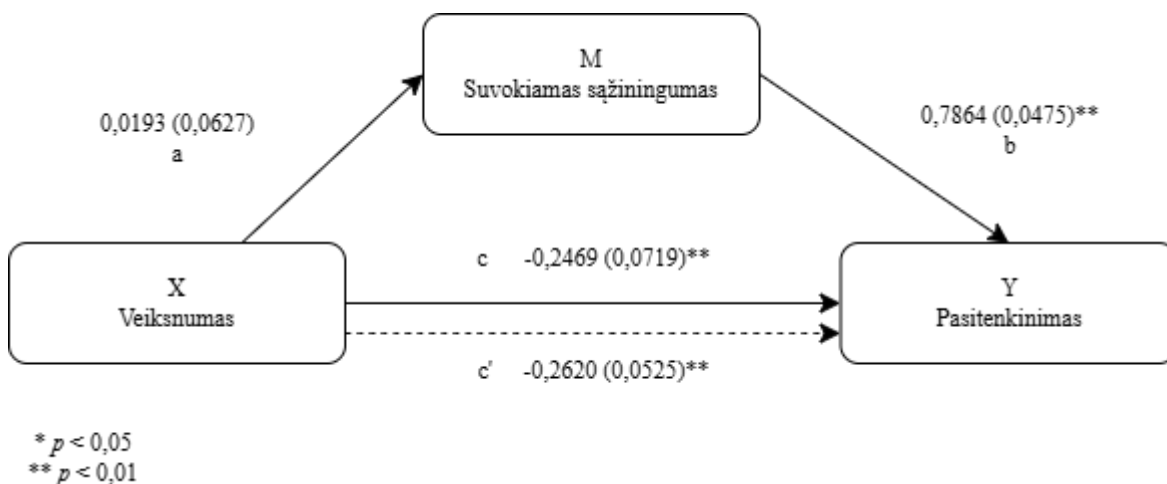
Tiesioginė, netiesioginė ir suminė veiksnų įtaka pasitenkinimui pateikiama žemiau (žr. **44 lentelė**).

44 lentelė. Tiesioginė, netiesioginė (per suvokiamą sąžiningumą) ir suminė veiksnų įtaka pasitenkinimui (N=313)

Kelias	Efektas (EF)	95% pasikliautinis intervalas*	
		LLCI	ULCI
Tiesioginė polinkio priskirti veiksnų įtaka pasitenkinimui			
$V \rightarrow P (c')$	-0,2620	-0,3653	-0,1588
Netiesioginė polinkio priskirti veiksnų įtaka pasitenkinimui			
$V \rightarrow S \rightarrow P (a*b)$	0,0152	-0,0873	0,1169
Suminė polinkio priskirti veiksnų įtaka pasitenkinimui			
$V \rightarrow P (c)$	-0,2469	-0,3884	-0,1054

Analizės rezultatai rodo, kad egzistuoja tiesioginis, neigiamas polinkio priskirti veiksnų poveikis pasitenkinimui (EF = -0,2469, LLCI = -0,3884; ULCI = -0,1054), mažėjant suvokiamam veiksnui,

didėja pasitenkinimas ir atvirkščiai. Netiesioginis efektas, medijuojant suvokiamam sąžiningumui, nėra reikšmingas (EF = 0,0152; LLCI = -0,0873; ULCI = 0,1169). Iliustruoti rezultatai pateikti 9 pav..



9 pav. Veiksnumo įtaka pasitenkinimui paslauga per suvokiamą sąžiningumą

Remiantis gautais rezultatais, hipotezė **H7b**: *Suvokiamas sąžiningumas medijuoja ryšį tarp veiksnio ir pasitenkinimo taip, kad: veiksnys daro teigiamą poveikį sąžiningumui, kuris savo ruožtu teigiamai veikia pasitenkinimą, nepatvirtinama.*

Apibendrinant tyrimo duomenis, pristatomi hipotezių patikros rezultatai (žr. 45 lentelė).

45 lentelė. Empirinio tyrimo hipotezių tikrinimo rezultatai

Hipotezė	Rezultatas	Pagrindimas
H1a : DI priimtas sprendimas (palyginus su žmogaus sprendimu) sukels polinkį priskirti didesnę moralumą sprendimo priėmėjui	Nepatvirtinama	Žr. 29 lentelė, 30 lentelė ir 31 lentelė
H1b : Žmogaus priimtas sprendimas (palyginus su DI sprendimu) sukels polinkį priskirti didesnę veiksnumą sprendimo priėmėjui	Patvirtinama	Žr. 29 lentelė, 32 lentelė ir 33 lentelė
H2a : Palankus sprendimas (palyginus su nepalankiu) sukels polinkį priskirti didesnę moralumą sprendimo priėmėjui	Patvirtinama	Žr. 29 lentelė, 30 lentelė ir 31 lentelė
H2b : Palankus sprendimas (palyginus su nepalankiu) sukels polinkį priskirti didesnę veiksnumą sprendimo priėmėjui	Patvirtinama	Žr. 29 lentelė, 32 lentelė ir 33 lentelė
H3a : Kai sprendimas yra palankus, nepriklausomai nuo to, ar sprendimą priėmė DI, ar žmogus, sprendimo priėmėjas bus suvokiamas, kaip labiau moralus	Nepatvirtinama	Žr. 29 lentelė, 30 lentelė ir 31 lentelė
H3b : Kai sprendimas yra nepalankus, DI sprendimų priėmėjas (palyginus su sprendimų priėmėju žmogumi) bus suvokiamas, kaip labiau veiksnus	Nepatvirtinama	Žr. 29 lentelė, 32 lentelė ir 33 lentelė
H4a : Vartotojų polinkis priskirti moralumą sprendimo priėmėjui daro teigiamą poveikį suvokiamam sąžiningumui	Patvirtinama	Žr. 34 lentelė ir 35 lentelė
H4b : Vartotojų polinkis priskirti veiksnumą sprendimo priėmėjui daro teigiamą poveikį suvokiamam sąžiningumui	Patvirtinama	Žr. 34 lentelė ir 35 lentelė
H5a : Sprendimo priėmėjas moderoja ryšį tarp polinkio priskirti moralumą ir sąžiningumo: kai sprendimą priima žmogus, ryšys tarp moralumo ir sąžiningumo bus stipresnis palyginus su situacija, kai sprendimą priima DI	Nepatvirtinama	Žr. 36 lentelė

H5b: Sprendimo priėmėjas moderoja ryšį tarp polinkio priskirti veiksnumą ir sąžiningumo: kai sprendimą priima DI, ryšys tarp veiksnumo ir sąžiningumo bus stipresnis palyginus su situacija, kai sprendimą priima žmogus	Nepatvirtinama	Žr. 37 lentelė
H6: Suvokiamas sprendimo priėmėjo sąžiningumas teigiamai veikia vartotojų pasitenkinimą paslauga	Patvirtinama	Žr. 38 lentelė ir 39 lentelė
H6a: Sprendimo priėmėjas moderoja ryšį tarp sąžiningumo ir pasitenkinimo taip, kad: kai sprendimą priima DI, ryšys tarp sąžiningumo ir pasitenkinimo bus stipresnis, palyginus su situacija, kai sprendimą priima žmogus	Nepatvirtinama	Žr. 40 lentelė
H7a: Suvokiamas sąžiningumas medijuoja ryšį tarp moralumo ir pasitenkinimo taip, kad: moralumas daro teigiamą poveikį sąžiningumui, kuris savo ruožtu teigiamai veikia pasitenkinimą	Patvirtinama	Žr. 41 lentelė ir 42 lentelė
H7b: Suvokiamas sąžiningumas medijuoja ryšį tarp veiksnumo ir pasitenkinimo taip, kad: veiksnumas daro teigiamą poveikį sąžiningumui, kuris savo ruožtu teigiamai veikia pasitenkinimą	Nepatvirtinama	Žr. 43 lentelė ir 44 lentelė

Tyrimas atskleidė, kad moraliai abejotinosiose paslaugų teikimo aplinkybėse, nepriklausomai, ar sprendimą priima žmogus, ar DI, sprendimų priėmėjų gebėjimas elgtis moraliai vertinamas panašiai. Polinkis priskirti veiksnumą yra didesnis, kai sprendimą priima žmogus. Polinkiui priskirti moralinį veiksnumą reikšmingą poveikį turi sprendimo palankumas. Kai sprendimas palankus, vartotojai yra linkę priskirti didesnę moralumą abiem sprendimo priėmėjams – DI ir žmogui. Nepalankaus sprendimo atveju, suvokiamas sprendimo moralumas reikšmingai sumažėja. Šiam poveikiui sprendimo priėmėjo tipas įtakos neturi. Polinkis priskirti veiksnumą yra veikiamas tiek sprendimo palankumo, tiek sprendimo priėmėjo. Kai sprendimas palankus, polinkis priskirti veiksnumą yra didesnis sprendimo priėmėjo žmogaus atveju nei DI. Kuomet sprendimas nepalankus, sprendimo priėmėjai yra vertinami, kaip mažiau veiksnūs, pasižymintys didesne priklausomybe nuo kitų. Vis dėlto, šis efektas yra silpnas. Suvokiamas DI veiksnumas palankaus ir nepalankaus sprendimo situacijoje išlieka panašus. Polinkis priskirti moralinį veiksnumą turi teigiamą poveikį suvokiamam sąžiningumui. Šiam ryšiui sprendimo priėmėjo poveikis įtakos neturi. Suvokiamas sprendimo priėmėjo sąžiningumas teigiamai veikia pasitenkinimą, bet sprendimo priėmėjo tipas šio poveikio nemoderoja. Nustatyta, kad suvokiamas sąžiningumas medijuoja ryšį tarp polinkio priskirti moralumą ir pasitenkinimo, tuo tarpu veiksnumo atveju šis ryšys nėra reikšmingas. Nepaisant to, kad hipotezė apie tiesioginę polinkio priskirti moralinį veiksnumą įtaka vartotojų pasitenkinimui kelta nebuvo, pastebėtas statistiškai reikšmingas ryšys.

4.5. Vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis empirinio tyrimo rezultatų diskusija

Šiuo baigiamuoju magistro projektu siekta teoriškai ir empiriškai pagrįsti vartotojų polinkio priskirti moralinį veiksnumą DI, ir žmogaus sprendimams poveikį pasitenkinimui paslaugomis. Mokslinės literatūros analizė atskleidė, jog egzistuoja skirtumai tarp to, kaip vartotojai vertina žmogaus ir DI sprendimus paslaugose. Vis dėlto, polinkio priskirti moralinį veiksnumą reikšminys rinkodaros ir susijusiose mokslo srityse nagrinėtas nedaug, nepaisant galimo poveikio vartotojų atsakui. Plačiausiai analizuotos atskiros šio reiškinio dedamosios, o mokslininkų išvados nėra vieningos.

Nors autoriai sutaria, kad moralinis veiksnumas yra žmonėms būdingas bruožas, empiriniai tyrimai grindžia, jog socialinėje sąveikoje šios savybės priskiriamos ir DI. Dalis mokslininkų bendruomenės narių teigia, jog moraliniu požiūriu svarbiose situacijose, moralinės savybės labiau priskiriamos DI,

kiti – žmogui, arba abiem sprendimų priėmėjams panašiai. Dauguma mokslinių darbų atskleidžia, jog vartotojai yra linkę priskirti didesnę veiksnumą žmonėms. Kita vertus, egzistuoja atveju, kai paslaugose vartotojai sprendimus priimančius žmones sutapatina su sisteminiu subjektu ir nėra linkę priskirti veiksnumo (Söderlund, 2024). Žinoma, jog moralinis veiksnumas turi įtakos žmonių atsakui, dažniausiai atskleista teigiama įtaka pasitikėjimui. Tačiau, iki šiol atlikti tyrimai nepateikė vieningo atsakymo, kuriam sprendimo priėmėjui, žmogui ar DI, vartotojai bus linkę priskirti didesnę moralinę veiksnumą, ir kaip tai veiks šiame kontekste dar netirtą vartotojų atsako formą, pasitenkinimą. Ayad'as ir Plaks'as (2025) siūlo atlikti tolimesnius tyrimus ir įvertinti, kaip polinkis priskirti moralinę veiksnumą DI ir žmogui, pasireiškia, kuomet abu sprendimo priėmėjai priima identiškus, moraliai dviprasmiškus sprendimus. Šis magistro projektas atliepia minėtų mokslininkų skatinimus ir analizuoja, kaip vartotojų polinkis priskirti moralinę veiksnumą DI ir žmogaus sprendimams veikia pasitenkinimą paslauga, moraliai abejotinoje paslaugų aplinkybėje.

Literatūros analizė atskleidė, jog vartotojų suvokimas apie moralinę veiksnumą gali būti lemiamas daugelio veiksnių, kaip sąveikos tipas ar kontekstinė aplinka. Pastebėta, jog šis reiškinys menkai nagrinėtas atsižvelgiant į sprendimo palankumo poveikį. Empirinės išvados atskleidžia, jog sprendimo palankumas turi įtakos suvokiamam paslaugų subjekto vertinimui ir veiksnumui (Söderlund, 2024), todėl daroma prielaida, jog šis veiksnys gali turėti įtakos ir suvokiamam moralumui. Garvey'is ir kt. (2023) prisideda prie stebimos sprendimo palankumo svarbos tiriamam reiškiniui, rekomenduodami praplėsti tyrimų sritį, nustatant, kaip vartotojų vertinimai apie skirtingus sprendimo priėmėjus skiriasi palankių ir nepalankių sprendimų atveju. Todėl, šiame projekte keliama tiesioginė sprendimo palankumo poveikio hipotezė ir moderuojanti, tarp polinkio priskirti moralinę veiksnumą DI ir žmogui.

Ankstesniuose tyrimuose dažniausiai analizuotas polinkio priskirti moralinę veiksnumą ryšys tarp suvokiamos sprendimo priėmėjo atsakomybės ar kaltės. Pastebėta, jog polinkio priskirti moralinę veiksnumą ir suvokiamo sąžiningumo ryšys dar nebuvo tirtas. Suvokiamas sąžiningumas siejamas su moraliniu veiksnumu, nes reikalauja gebėjimo situaciją vertinti moraliniu požiūriu ir priimti sprendimus, atsižvelgiant į galimas pasekmes. Todėl, keltos hipotezės apie teigiamą polinkio priskirti moralinę veiksnumą ir suvokiamo sąžiningumo ryšį. Ši prielaida grindžiama sąžiningumo teorija (Folger ir Cropanzano, 2001) bei procedūrinio ir tarpasmeninio sąžiningumo dimensijomis (Colquitt et al., 2001). Sąžiningumo teorija nurodo, jog žmonės yra linkę vertinti sprendimo priėmėjo atskaitomybę ir elgsenos atitikimą moralės normoms. Procedūrinis sąžiningumas siejamas su gebėjimu kontroliuoti situaciją, priimti autonomiškus sprendimus, tuo tarpu tarpasmeninis sąžiningumas nusako subjekto gebėjimą priimti moraliai teisingus sprendimus. Suvokiamas sąžiningumas yra glaudžiai susijęs su vartotojų atsaku į paslaugos sandorį, ypač vartotojų pasitenkinimo aspektu. Remiantis ankstesnėmis empirinėmis išvadomis apie teigiamą moralinio veiksnumo poveikį vartotojų atsakui ir sąžiningumo poveikį vartotojų pasitenkinimui, šiame baigiamajame magistro projekte, suvokiamas sprendimo priėmėjo sąžiningumas traktuotas kaip ryšys tarp polinkio priskirti moralinę veiksnumą ir pasitenkinimo paaiškinantis psichologinis mechanizmas.

Šio baigiamojo magistro projekto konstruktai – polinkis priskirti moralinę veiksnumą, sprendimo palankumas, suvokiamas sąžiningumas ir vartotojų pasitenkinimas – nagrinėti atskirai, tačiau jų tarpusavio sąveika dar nebuvo tirta. Atsižvelgiant į mokslinės literatūros spragas, šiuo projektu siekta praplėsti teorines ir empirines išvadas apie minėtų konstrukto ryšius. Remiantis mokslinės literatūros analizės rezultatais, iškelta 12 hipotezių. Siekiant patikrinti iškeltas hipotezes, taikyta kiekybinio tyrimo strategija, vykdant internetiniais scenarijais grįstą eksperimentą.

Tyrimo rezultatai rodo, jog moraliai abejotinosiose paslaugų teikimo aplinkybėse, vartotojai yra linkę priskirti moralumą abiem sprendimo priėmėjams, DI ir žmogui, panašiai. Nors vidurkių skirtumai atskleidžia, jog DI priskiriamas šiek tiek didesnis moralumas, statistiškai reikšmingo skirtumo nėra. Šie rezultatai prisideda prie Zhang ir kt. (2023) tyrimo išvadų, teigiančių, jog DI ir žmogaus sprendimai vertinami lygiavertiškai, kai moralinė situacija sukelia žalą. Panašios išvalgos atspindi ir Banks (2021) tyrimo rezultatuose, jog DI ir žmogaus elgsena moraliniu požiūriu vertinama panašiai. Tačiau, polinkis priskirti veiksnumą yra didesnis sprendimo priėmėjui žmogui, lyginant su sprendimo priėmėju DI. Šie rezultatai prisideda prie esamų išvadų apie didesnę polinkį priskirti veiksnumą žmonėms nei DI (Pitardi et al., 2022; Pavone et al., 2023; Arikan et al., 2023; Swiderska & Küster, 2021). Be to, praplečia rinkodaros ir vartotojų elgsenos tyrimų lauką, nes pateikia naujus, iki šiol netyrinėtus aspektus, apie egzistuojančius skirtumus tarp vartotojų polinkio priskirti moralinį veiksnumą paslaugą teikiančiam DI ir žmogui.

Nustatytas statistiškai reikšmingas sprendimo palankumo poveikis polinkiui priskirti moralinį veiksnumą. Sprendimo palankumas lemia vartotojų polinkį priskirti didesnę moralinį veiksnumą. Tačiau, šis ryšys nepriklauso nuo to, ar sprendimą priėmė DI, ar žmogus. Ši rezultatai atskleidžia, jog vartotojai, suvokdami sprendimų priėmėjų šališkumą, yra linkę pateisinti sprendimų rezultatą, kai iš jo gauna naudos. Gauti duomenys praturtina ankstesnių tyrimų rezultatus, jog sprendimo palankumas teigiamai veikia sprendimų priėmėjų vertinimus, nepriklausomai nuo jų tipo (Akter et al., 2022; Choi & Chao, 2024), net ir moraliniu požiūriu abejotinosiose paslaugų teikimo aplinkybėse. Kai sprendimas nepalankus, vartotojai yra linkę priskirti mažesnę moralinį veiksnumą. Mažesnis moralinis veiksnumas priskiriamas abiem sprendimo priėmėjams, dėl numanomo sumažėjusio autonomiškumo ir patiriamos tiesioginės žalos. Šis reiškinys gali būti aiškinamas atribucijos teorija (Jones & Davis, 1965), kuomet sprendimas nepalankus, vartotojai gali išsamiau analizuoti išorines neigiamo rezultato priežastis, pavyzdžiui sprendimo priėmėjo savybes, ir jas vertinti griežčiau. Priešingai nei buvo hipotetizuota, nepalankaus sprendimo atveju, žmogus, lyginant su DI, suvokiamas, kaip labiau moraliai veiksnus. Ši išvada iš dalies prieštarauja Choi ir Chao (2024), ir Söderlund'o (2024) tyrimų rezultatams, kurie teigia, jog nepalankūs DI sprendimai, palyginti su žmogumi, bus vertinami palankiau, tokiu būdu DI priskiriant teigiamus atributus ir didesnę veiksnumą. Gauti rezultatai siejasi su proto teorija (Premack & Woodruff, 1978), teigiančia, jog žmonės vertina kitų veiksnumą, atsižvelgdami į stebimo subjekto gebėjimą veikti sąmoningai, priimti racionalius sprendimus ir kontroliuoti situaciją. Nepaisant to, jog žmonės yra linkę priskirti veiksnumo savybes tiek žmogui, tiek DI, šis tyrimas prisideda prie ankstesnių išvadų, jog žmogus dažniau suvokiamas, kaip gebantis priimti moralesnius ir sąmoningesnius sprendimus (Thellman et al., 2022). Tyrimo rezultatai atskleidžia, jog palankaus ir nepalankaus sprendimo atveju, polinkis priskirti veiksnumą DI, išlieka žemesnis nei žmogui. Manoma, jog ši išvada gali būti pagrįsta vartotojų suvokiama didesne DI priklausomybe nuo kitų (pavyzdžiui, programuotojų). Atsižvelgiant į tai, jog sprendimo palankumo tema, lyginant sprendimo priėmėją DI ir žmogų, yra dar nauja (Yalcin et al., 2022), būtina atlikti tolimesnius tyrimus ir toliau analizuoti galimas šio reiškinio priežastis.

Tyrimas parodė, kad polinkis priskirti moralinį veiksnumą turi reikšmingą teigiamą poveikį suvokiamam sprendimo priėmėjo sąžiningumui. Atsižvelgiant į tirtą moraliai abejotiną paslaugų kontekstą, vartotojų priskiriamas mažesnis moralinis veiksnumas turėjo įtakos mažesniai suvokiamam sąžiningumui. Šie rezultatai atitinka teorines prielaidas, jog sąžiningumo teorija (Folger ir Cropanzano, 2001) ir procedūrinio bei tarpasmeninio sąžiningumo dimensijos (Colquitt et al., 2001) pagrindžia numanomas sąsajas apie ryšį tarp polinkio priskirti moralinį veiksnumą ir

suvokiamo sąžiningumo. Gauti rezultatai grindžia prielaidą, jog sprendimo priėmėjo sąžiningumas indikuoja apie tvirtai pagrįstus ketinimus elgtis laikantis moralinių ir etinių normų. Vis dėl to, nustatyta, jog sprendimo priėmėjas šio ryšio nemoderuoja. Tokiu būdu, šis tyrimas paneigia Wu ir kt. (2024), Qin'o ir kt. (2023), Helberger ir kt. (2020) rezultatus ir prisideda prie Noble, Foster & Craig (2021) ir Lee (2018) išvadų. Galima teigti, kad priimant moraliai abejotinus sprendimus paslaugose, skirtumo tarp DI ir žmogaus suvokiamo sąžiningumo, nėra.

Nustatytas statistiškai reikšmingas suvokiamo sąžiningumo poveikis vartotojų pasitenkinimui. Žvelgiant į tyrimo rezultatus, bendras šių tyrimo konstrukto vertinimas buvo labiau neigiamas. Todėl, moraliai abejotinosiose paslaugų situacijose, suvokiamas mažesnis sprendimo priėmėjo sąžiningumas turi įtakos mažesniai vartotojų pasitenkinimui paslauga. Sprendimo priėmėjo tipas – DI arba žmogus – nesąlygojo reikšmingų skirtumų ryšio pobūdžiui. Tokiai išvadai pagrindą gali suteikti procedūrinio teisingumo teorija (Colquitt et al., 2001), kuri nurodo, jog sąžiningumas siejamas su priimto sprendimo skaidrumu ir nešališkumu. Vadovaujantis šiuo požiūriu, moraliai abejotinosiose paslaugų aplinkybėse, vartotojų dėmesys gali labiau krypti į paties sprendimo rezultatą nei į subjektą, priėmusį sprendimą. Šio tyrimo išvados prisideda prie suvokiamo sąžiningumo ir vartotojų pasitenkinimo tyrimų (Yang & Lee, 2024; Koh et al., 2025, Christ-Brendemühl & Schaarschmidt, 2022) bei suteikia naujų žinių, apie skirtingų paslaugos teikėjų, žmogaus ir DI, sąžiningumo įtaką vartotojų pasitenkinimui. Be to, ankstesni empiriniai tyrimai atlikti restoranų (Ahmed et al., 2023), viešbučių ir turizmo (Alderighi et al., 2022), oro linijų (Setiawan et al., 2020) ar telekomunikacijų (Alzoubi et al., 2020) paslaugų kontekste, todėl šis tyrimas užpildo esamą spragą, leidžiant geriau suprasti, kaip šis reiškinys pasireiškia profesionaliose grožio paslaugose.

Kitas, šiam tyrimui labai svarbus ryšys tarp polinkio priskirti moralinį veiksnumą ir pasitenkinimo, medijuojant suvokiamam sąžiningumui, patvirtintas tik iš dalies. Duomenys atskleidė, kad suvokiamas sąžiningumas paaiškina dalį ryšio tarp polinkio priskirti moralumą ir pasitenkinimo. Nustatytas ir tiesioginis moralumo, ir pasitenkinimo ryšys, kuris buvo statistiškai stipresnis. Tuo tarpu, suvokiamo sąžiningumo tarpininkavimo poveikis tarp polinkio priskirti veiksnumą ir vartotojų pasitenkinimo paslauga, nepatvirtintas. Nors šio tyrimo ribose tiesioginis poveikis nebuvo prognozuojamas, analizės rezultatai parodė, jog tarp šių tyrimo konstrukto, veiksnumo ir pasitenkinimo, egzistuoja statistiškai reikšmingas, neigiamos krypties ryšys. Gauti rezultatai rodo, jog suvokiamas sprendimo priėmėjo veiksnumas daro įtaką suvokiamam sąžiningumui, suvokiamas sąžiningumas teigiamai veikia pasitenkinimą, tačiau veiksnumas, per sąžiningumą, nebūtinai formuoja vartotojų pasitenkinimą. Žinant, jog moralumas turėjo tiesioginį poveikį pasitenkinimui paslauga, keliama prielaida, jog polinkis priskirti moralinį veiksnumą vartotojų pasitenkinimą veikia tiesiogiai. Atsižvelgiant į tai, jog esama literatūra nepateikia įžvalgų apie šiuos ryšius, ši išvada reikalauja gilesnio teorinio pagrindimo ir empirinių išvadų.

Toliau pateikiami vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio pasitenkinimui paslaugomis **tyrimo apribojimai**:

1. Tyrimo imčiai pasiekti, taikytas netikimybinis patogusis atrankos metodas, todėl gauti rezultatai neapibūdina populiacijos visumos. Remiantis statistiniu ir nestatistiniu tyrimo imties nustatymu, minimali riba pasiekta, tačiau siekiant didesnio duomenų patikimumo, turi būti užtikrinama didesnė tyrimo imties aprėptis.
2. Internetiniais scenarijais grįstas eksperimentinis tyrimo dizainas neleidžia pilnai kontroliuoti tyrimo aplinkos sąlygų. Nepaisant galimybės greitai pasiekti didelę tyrimo imtį, egzistuoja

išorinių, su tyrimų nesusijusių veiksmų įtakos tikimybė, galinti paveikti tyrimo dalyvių atsakymų kokybę. Siekiant didesnio duomenų tikslumo, panašus tyrimas galėtų būti atliktas realioje aplinkoje, įvertinant kitus svarbius veiksmus, kaip respondentų emocinė reakcija.

3. Remiantis literatūroje pateikiamomis išvadomis, daugelis tyrimo konstrukto gali būti veikiami kultūrinių ir kontekstinių aplinkybių. Nors siekta apimti geografiškai įvairių tyrimo imčių, respondentų pasiskirstymas nebuvo vienodas. Todėl, būtina atlikti tolesnius tyrimus bei reprezentuoti įvairius kultūrinius ir kontekstinius skirtumus, kurie gali turėti įtakos tyrimo rezultatams.
4. Tyrimo kontekstui pasirinktos profesionalios grožio paslaugos, todėl, gauti rezultatai negali būti apibendrinami ir taikytini kitiems verslo sektoriams. Dėl šios priežasties, svarbu įvertinti, kaip tyrimo konstrukto ryšiai atspindi kitose industrijose.
5. Moraliai abejotinos situacijos suintensyvėjimui, išryškintas lyčių šališkumas, tačiau žinoma, jog lyčių priešprieša visuomenėje yra mažėjanti, todėl tai galėjo lemti mažesnę tyrimo dalyvių jautrumą ir tyrimo rezultatus.

Tolesnės tyrimų kryptys:

1. Pravartu įvertinti daugiau polinkio priskirti moralinį veiksmą raiškos sąlygų. Polinkis priskirti moralinį veiksmą gali būti vertinamas daugiau nei dvejais lygiais (žemas ar aukštas). Galima išvesti keturis, pavyzdžiui: aukštas moralumas / žemas veiksmas, aukštas moralumas/ aukštas veiksmas/, žemas moralumas / žemas veiksmas, žemas moralumas / aukštas veiksmas.
2. Įtraukus vizualinį sprendimo priėmėjo (DI ir žmogaus) atvaizdą, gali išryškėti papildomos įžvalgos. DI antropomorfizmas gali turėti įtakos polinkiui priskirti moralinį veiksmą (Wester et al., 2024). Todėl, būsimuose tyrimuose gali būti vertinama fizinio įkūnijimo bei antropomorfizmo įtaka šiame tyrime nustatytiems ryšiams.
3. Aktualu įvertinti, kaip kinta polinkis priskirti moralinį veiksmą tiesioginėje vartotojų, paslaugą teikiančių DI sistemų ir žmonių sąveikoje.
4. Nustatytas teigiamas sprendimo palankumo poveikis polinkiui priskirti moralinį veiksmą, svarbu įvertinti kitus veiksmus, pavyzdžiui, pasitikėjimo poveikį.
5. Šio tyrimo kontekstas susijęs su moraliai abejotina paslaugų teikimo aplinka, nagrinėjant šią temą toliau, tikslinga įvertinti skirtumus, esant moraliniu požiūriu teigiamai situacijai. Nustatyti, kaip keičiasi vartotojų suvokimas, kai priimamas moraliai teisingas sprendimas.
6. Rekomenduojama analizuoti šio reiškinio poveikį kitoms vartotojų atsako formoms, pavyzdžiui, elgsenos ketinimams. Söderlund'as (2024) savo tyrimu įrodė, jog suvokiamas subjekto veiksmas teigiamai veikia komunikaciją iš lūpų į lūpas. Ateities tyrimuose galėtų būti nagrinėjama, kaip polinkis priskirti moralinį veiksmą veikia vartotojų ketinimus rekomenduoti paslaugą kitiems (angl. *word-of-mouth*).

Išvados

Teoriškai ir empiriškai patikrinus vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikį pasitenkinimui paslaugomis, formuluojamos šios išvados:

1. Polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams reiškinyje nagrinėtas psichologijos, sociologijos, žmogaus ir kompiuterio sąveikos moksliniuose tyrimuose, kur nustatytas reikšmingas poveikis žmonių pasitikėjimui ir kitoms psichologinėms, ir emocinėms reakcijoms. Nepaisant sparčios DI plėtros versle ir įvairiuose paslaugų sektoriuose, kuomet vartotojų ir žmogaus sąveiką vis dažniau keičia DI sistemos, šis reiškinyje rinkodaros ir susijusiose mokslo kryptyse tirtas siaurai. Plačiausiai tyrinėti atskiri vartotojų polinkio priskirti moralinį veiksnumą aspektai, kaip suvokiamas moralumas, polinkis priskirti veiksnumą, atsakomybę ar kaltę sprendimų priėmėjui DI ir žmogui, moraliai svarbiuose rinkodaros ar paslaugų kontekstuose. Pastebėta, jog šioje tematikoje vyrauja skirtingos mokslinių tyrimų išvados. Dalis autorių teigia, kad žmonės DI nėra linkę priskirti žmogui būdingų moralinio veiksnumo savybių, kiti pateikia priešingas išvadas. Be to, nesutariama, kaip mažesnis ar didesnis polinkis priskirti moralinį veiksnumą veikia vartotojų reakciją į sprendimo priėmimą ir paslaugą plačiąja prasme. Identifikuotas mokslinių tyrimų ribotumas, analizuojant moraliai dviprasmiškas paslaugų teikimo situacijas, esant vartotojų atžvilgiu palankiam ir nepalankiam sprendimui (Ayad & Plaks, 2025, Garvey et al., 2023). Trūksta žinių, kaip minėtas reiškinyje siejasi su vartotojų suvokiamu DI ir žmogaus sąžiningumu ir, kaip tai veikia vartotojų atsaką, pavyzdžiui, pasitenkinimą (Garvey et al., 2023). Naujausi tyrimai atskleidžia, kad polinkis priskirti moralinį veiksnumą turi įtakos vartotojų pasitikėjimui ir paslaugos priimtinumui, tačiau poveikis pasitenkinimui dar netirtas. Todėl, nustatytas poreikis ištirti, kaip vartotojų polinkis priskirti moralinį veiksnumą DI ir žmogaus sprendimams veikia pasitenkinimą paslauga.
2. Mokslinėje literatūroje, moralinis veiksnumas kvalifikuojamas subjekto laisva valia, autonomiškumu, sąmoningumu bei gebėjimu elgtis pagal moralės normas. Akcentuojama, jog šis reiškinyje nebėra tik žmogų nusakantis bruožas, gebėjimas veikti racionaliai moraliai, šiandien būdingas ir DI. Šios DI moralinio veiksnumo sąlygos apibūdinamos per dirbtinių moralinių agentų (DMA) sąvoką, kuri nurodo bet kokio pobūdžio DI technologijos gebėjimą sąmoningai priimti moralinius sprendimus. Nors etikos mokslo šalininkai dviprasmiškai vertina lygiavertišką žmogaus ir DI moralinio veiksnumo vertinimą, sociologijos ir žmogaus-kompiuterio sąveikos mokslininkai pabrėžia, jog žmonės, dalyvaudami socialinėje sąveikoje su DI, yra linkę priskirti tam tikro lygio moralinį veiksnumą. Polinkis priskirti moralinį veiksnumą aiškinamas suvokiamo moralinio veiksnumo (SMV) reiškiniu (Banks, 2019). Suvokiamas moralinis veiksnumas nusako žmonių polinkį priskirti moralę ir veiksnumą bet kokios rūšies stebimam subjektui (gyvūnui, žmogui ar DI), todėl šis reiškinyje aiškinamas per dvi minėtas sąlygas: moralumą ir veiksnumą. Polinkis priskirti moralę, grindžiamas plačiai pripažinta, 2004 m. Haidt'o ir Joseph'o pasiūlyta moralinių pagrindų teorija (MPT), kuri nurodo, jog stebimi subjektai vertinami, remiantis penkiais moraliniais principais: rūpestis / žala, sąžiningumas / apgavystė, lojalumas / ištikimybė, autoritetas / subversija, tyrumas / nuopuolis. Literatūra akcentuoja, jog remiantis šiais moralės pagrindais, vertinamas subjektų elgsenos atitikimas moralės normoms. Tuo tarpu 1978 m., Premack'o ir Woodruff'o pasiūlyta proto teorija (PT) paaiškina polinkį priskirti veiksnumą. Teorija nurodo, kad žmonės vertina stebimo subjekto psichinę būseną, ją lygina su savuoju proto modeliu ir prognozuoja galimus subjekto gebėjimus. Suvokiamas protas nusako stebimo subjekto veiksnumo ir patirties prielaidą, kurios gali būti priskiriamos kartu arba atskirai. Veiksnumas

apibūdina suvokiamą subjekto savikontrolę, gebėjimą racionaliai, sąmoningai priimti sprendimus. Todėl, šiame baigiamajame magistro projekte, remiantis MPT ir PT, paaiškinamas vartotojų polinkis priskirti moralinį veiksnumą sprendimų priėmėjui DI ir žmogui. Empiriniai duomenys atskleidžia, jog polinkis priskirti moralinį veiksnumą DI ir žmogui, priklauso nuo daugelio susijusių veiksnių, kaip psichologiniai mechanizmai, sąveikos pobūdis ar situacijos kontekstas. Paslaugose, vartotojų ir paslaugą teikiančių subjektų sąveika gali stipriai priklausyti nuo sprendimo palankumo. Tyrimai atskleidžia, jog sprendimo palankumas turi įtakos vartotojų reakcijoms, ypač nepageidaujamo rezultato atveju. Nustatyta, jog egzistuoja sąsajos tarp nepalankaus sprendimo ir sumažėjusio polinkio priskirti veiksnumą. Šį reiškinį paaiškina Jones'o ir Davis'o, 1965 m. atribucijos teorija, kuomet neigiamai traktuojama situacija skatina ieškoti išorinių priežasčių, lėmusių nepageidaujamą rezultatą ir jį vertinti griežčiau. Mokslinės literatūros analizė atskleidė, jog polinkis priskirti moralinį veiksnumą gali turėti įtakos suvokiamam DI ir žmogaus sąžiningumui. Vadovaujantis Folger'io ir Cropanzano, 2001 m. pasiūlyta sąžiningumo teorija, suvokiamas sąžiningumas apibūdina stebimo subjekto atskaitomybės vertinimą, atsižvelgiant į žalos mastą, elgsenos modelį ir moralės normų laikymąsi. Remiantis Colquitt'o ir kt. (2001) procedūrinio ir tarpasmeninio sąžiningumo dimensijomis, vartotojai vertina priimto DI ir žmogaus gebėjimą kontroliuoti situaciją ir priimti moraliai teisingus sprendimus, gerbiant vartotojų orumą ir tinkamai reflektuojant į sprendimo rezultatą. Šios savybės apibūdina glaudžias sąsajas su moraliniu veiksnumu. Moksliniai tyrimai atskleidžia, jog suvokiamas sąžiningumas paslaugose, teigiamai veikia vartotojų atsaką, skatina lojalumą, komunikaciją iš lūpų į lūpas, pasitikėjimą bei pasitenkinimą. Tačiau, empiriniuose darbuose dar tik pradėta nagrinėti, kokią įtaką vartotojų pasitenkinimui turi suvokiamas skirtingų sprendimų priėmėjų, DI ir žmogaus, sąžiningumas. Be to, nėra žinoma, kaip suvokiamas sąžiningumas medijuoja ryšį tarp polinkio priskirti moralinį veiksnumą ir pasitenkinimo paslaugomis. Remiantis teorinėmis prielaidomis ir empirinėmis išvadomis, sudarytas conceptualus modelis, pagrindžiantis vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikį pasitenkinimui paslaugomis, įvertinant sprendimo palankumo ir suvokiamo sąžiningumo vaidmenį.

3. Ryšiams tarp konstrukto nustatyti ir hipotezėms patikrinti, taikyta kiekybinio tyrimo strategija, atliekant internetiniais scenarijais grįstą eksperimentą. Tyrime taikytas tarpgrupinis eksperimento dizainas, 2 (sprendimo priėmėjas: DI arba žmogus) $\times$ 2 (sprendimo palankumas: palankus arba nepalankus). Tyrimo duomenims rinkti pasitelktas apklausos metodas. Konstrukto – polinkio priskirti moralinį veiksnumą, suvokiamo sąžiningumo ir pasitenkinimo – matavimui naudotos empiriškai patikrintos matavimo skalės. Atliktas išankstinis skalių testavimas atliekant sintetinį tyrimą. Tyrimo imtis apskaičiuota nestatistinio lyginamojo ir statistinio A-priori imties dydžio nustatymo metodais. Faktinę tyrimo imtį, po dėmesio patikrinimo, sudarė 313 respondentai. Dėl turimų išteklių, respondentų atrankai taikytas netikimybinis patogusis atrankos metodas. Anketa talpinta internetinėje apklausų kūrimo platformoje *Qualtrics* ir viešinta asmeniais komunikacijos kanalais bei socialiniuose tinkluose. Statistinis tyrimo duomenų apdorojimas vykdytas naudojant *IBM SPSS* programinę įrangą.
4. Empirinis tyrimas atskleidė, jog moraliai abejotiname paslaugų kontekste, vartotojų polinkis priskirti moralumą buvo panašus abiejų sprendimo priėmėjų, DI ir žmogaus, atveju. Patvirtinta, jog didesnis veiksnumas priskirtas žmogui. Sprendimo palankumas skatina vartotojų polinkį priskirti moralinį veiksnumą sprendimų priėmėjui. Tačiau, tiek palankaus, tiek nepalankaus sprendimo atveju, polinkis priskirti moralinį veiksnumą nepriklausė nuo to, ar sprendimą priėmė DI, ar žmogus. Rezultatai patvirtino, jog vartotojų polinkis priskirti moralinį veiksnumą teigiamai

veikia suvokiamą sprendimų priėmėjų sąžiningumą. Vis dėlto, šis ryšys išlieka nepriklausomas nuo sprendimo priėmėjo tipo – DI arba žmogaus. Suvokiamas sprendimo priėmėjo sąžiningumas teigiamai veikia vartotojų pasitenkinimą, bet tas, kas priėmė sprendimą – DI arba žmogus – šio ryšio nemoderuoja. Nustatyta, kad suvokiamas sąžiningumas paaiškina ryšį tarp vartotojų polinkio priskirti moralumą ir pasitenkinimo paslauga. Nors hipotezė kelta nebuvo, pastebėta, jog šis, teigiamas ryšys, yra stipresnis tiesioginio poveikio atveju. Nepatvirtinta, jog suvokiamas sąžiningumas medijuotų polinkio priskirti veiksnumą poveikį pasitenkinimui paslauga, tačiau, rezultatai atskleidė, jog tiesioginis, neigiamos krypties ryšys tarp polinkio priskirti veiksnumą ir pasitenkinimo egzistuoja.

5. Atlikus vartotojų polinkio priskirti moralinį veiksnumą DI ir žmogaus sprendimams poveikio vartotojų pasitenkinimui paslaugomis tyrimą, pateikiamos **praktinės rekomendacijos ir tolimesnių tyrimų kryptys:**

- Remiantis tuo, jog moraliai dviprasmiški žmogaus ir DI sprendimai vertinami panašiai, organizacijos turėtų užtikrinti etinį reguliavimą, nepriklausomai nuo paslaugų teikėjo tipo. Diskriminacijos atveju, vartotojai yra linkę priskirti mažesnę moralumą DI ne mažiau nei žmonėms, todėl svarbu taikyti vienodus atsakomybės standartus ir atitinkamai reguliuoti DI sistemų veikimą.
- Atsižvelgiant į tai, jog vartotojai yra linkę priskirti didesnę veiksnumą paslaugą teikiančiam žmogui, sąveikaujant su DI, turi būti užtikrinama galimybė atsisakyti DI sprendimų ir, esant poreikiui, pratęsti komunikaciją su žmogumi ar kitu atsakingu subjektu.
- Kai sprendimas palankus, padidėja vartotojų polinkis priskirti moralinį veiksnumą, nepaisant šališko DI ar žmogaus sprendimo. Tai rodo, jog dalis vartotojų yra linkę iš dalies ignoruoti ir pasyviau reaguoti į pastebimą diskriminaciją, kai gauna geresnį pasiūlymą. Todėl, organizacijos turėtų turėti stiprią vertybinę poziciją, kuri atspindėtų rinkodaros komunikacijoje, siekiant edukuoti vartotojus, skatinti jų refleksiją ir lojalumą per vertybinius pagrindus. Ši informacija gali būti ypač aktuali toms vartotojų grupėms, kurios vertina organizacijos skaidrumą.
- Esant moraliai abejotinioms paslaugų teikimo aplinkybėms, priskiriamas mažesnis moralinis veiksnumas, tai veikia suvokiamą sprendimo priėmėjo sąžiningumą, tiek žmogaus, tiek DI atžvilgiu, ir sumažėjusį vartotojų pasitenkinimą. Šios išvados grindžia, jog šiuolaikiniame vartotojui nebėra svarbi tik paslauga, svarbus ir organizacijos vertybinis pamatas. Būtent dėl šių priežasčių, rinkodaros specialistai, turėtų aiškiai reprezentuoti moralinėmis ir etinėmis vertybėmis grįstą, ir sąžiningumu vedamą organizacijos poziciją.
- Tolesniuose tyrimuose rekomenduojama atkreipti dėmesį į kitus psichologinius mechanizmus, galinčius turėti įtaką vartotojų polinkiui priskirti moralinį veiksnumą DI ir žmogaus sprendimams. Įvertinti antropomorfizmo poveikį šiam reiškiniui. Išanalizuoti, kaip šis reiškinys kinta moraliniu požiūriu teigiamose paslaugų aplinkybėse. Nustatyti galimą poveikį kitoms vartotojų atsako formoms, pavyzdžiui ketinimams rekomenduoti. Daugiau rekomendacijų tolimesniems tyrimams pateikiama 4.5 darbo poskyryje.

Literatūros sąrašas

1. Ahmed, S., Al Asheq, A., Ahmed, E., Chowdhury, U. Y., Sufi, T., & Mostofa, Md. G. (2023). The intricate relationships of consumers' loyalty and their perceptions of service quality, price and satisfaction in restaurant service. *TQM Journal*, 35(2), 519–539, [žiūrėta 2025-03-09]. Prieiga per internetą: <https://doi.org/10.1108/TQM-06-2021-0158>
2. Ahn, J., Kim, J., & Sung, Y. (2024). The role of perceived freewill in crises of human-AI interaction: the mediating role of ethical responsibility of AI. *International Journal of Advertising*, 43(5), 847–873, [žiūrėta 2024-11-15]. Prieiga per internetą: <https://doi.org/10.1080/02650487.2023.2299563>
3. Ayad, R., & Plaks, J. E. (2025). Attributions of intent and moral responsibility to AI agents. *Computers in Human Behavior: Artificial Humans*, 3, 100107, [žiūrėta 2024-12-19]. Prieiga per internetą: <https://doi.org/10.1016/j.chbah.2024.100107>
4. Akter, S., Dwivedi, Y. K., Sajib, S., Biswas, K., Bandara, R. J., & Michael, K. (2022). Algorithmic bias in machine learning-based marketing models. *Journal of Business Research*, 144, 201–216, [žiūrėta 2025-02-23]. Prieiga per internetą: <https://doi.org/10.1016/j.jbusres.2022.01.083>
5. Akter, S., Hossain, M. A., Sajib, S., Sultana, S., Rahman, M., Vrontis, D., & McCarthy, G. (2023). A framework for AI-powered service innovation capability: Review and agenda for future research. *Technovation*, 125, 102768, [žiūrėta 2024-12-17]. Prieiga per internetą: <https://doi.org/10.1016/j.technovation.2023.102768>
6. Alderighi, M., Nava, C. R., Calabrese, M., Christille, J.-M., & Salvemini, C. B. (2022). Consumer perception of price fairness and dynamic pricing: Evidence from Booking.com. *Journal of Business Research*, 145, 769–783, [žiūrėta 2025-03-09]. Prieiga per internetą: <https://doi.org/10.1016/j.jbusres.2022.03.017>
7. Alordiah, C. O., & Chenube, O. (2023). Item Rotation in Scale Development: Exploring Principles and Strategies for Improving Scale Validity and Interoperability through Factor Analysis. *Nigerian Association of Social Psychologists*, 6 (2), 65-82, [žiūrėta 2025-05-01]. Prieiga per internetą: <https://www.nigerianjsp.com/index.php/NJSP/article/view/100>
8. Alzoubi, H., Alshurideh, M., Al Kurdi, B., & Inairat, M. (2020). Do perceived service value, quality, price fairness, and service recovery shape customer satisfaction and delight? A practical study in the service telecommunication context. *Uncertain Supply Chain Management*, 8(2), 579–588, [žiūrėta 2025-03-09]. Prieiga per internetą: <https://doi.org/10.5267/j.uscm.2020.2.005>
9. Angerschmid, A., Zhou, J., Theuermann, K., Chen, F., & Holzinger, A. (2022). Fairness and Explanation in AI-Informed Decision Making. *Machine Learning and Knowledge Extraction*, 4(2), 556–579, [žiūrėta 2025-03-03]. Prieiga per internetą: <https://doi.org/10.3390/make4020026>
10. Antony, J., Viles, E., Torres, A. F., Paula, T. I. de, Fernandes, M. M., & Cudney, E. A. (2020). Design of experiments in the service industry: a critical literature review and future research directions. *TQM Journal*, 32(6), 1159–1175, [žiūrėta 2024-11-19]. Prieiga per internetą: <https://doi.org/10.1108/TQM-02-2020-0026>
11. Araujo, T., Helberger, N., Kruijemeier, S., & de Vreese, C. H. (2020). In AI we trust? Perceptions about automated decision-making by artificial intelligence. *AI & Society*, 35(3), 611–623, [žiūrėta 2025-03-02]. Prieiga per internetą: <https://doi.org/10.1007/s00146-019-00931-w>
12. Arian, E., Altinigne, N., Kuzgun, E., & Okan, M. (2023). May robots be held responsible for service failure and recovery? The role of robot service provider agents' human-likeness. *Journal*

- of Retailing and Consumer Services, 70, 103175, [žiūrēta 2024-11-13]. Prieiga per internetą: <https://doi.org/10.1016/j.jretconser.2022.103175>
13. Bankins, S., Formosa, P., Griep, Y., & Richards, D. (2022). AI Decision Making with Dignity? Contrasting Workers' Justice Perceptions of Human and AI Decision Making in a Human Resource Management Context. *Information Systems Frontiers*, 24(3), 857–875, [žiūrēta 2025-03-02]. Prieiga per internetą: <https://doi.org/10.1007/s10796-021-10223-8>
 14. Banks, J. (2019). A perceived moral agency scale: Development and validation of a metric for humans and social machines. *Computers in Human Behavior*, 90, 363–371, [žiūrēta 2025-12-04]. Prieiga per internetą: <https://doi.org/10.1016/j.chb.2018.08.028>
 15. Banks, J. (2020). Theory of Mind in Social Robots: Replication of Five Established Human Tests. *International Journal of Social Robotics*, 12(2), 403–414, [žiūrēta 2025-12-04]. Prieiga per internetą: <https://doi.org/10.1007/s12369-019-00588-x>
 16. Banks, J. (2021). Good Robots, Bad Robots: Morally Valenced Behavior Effects on Perceived Mind, Morality, and Trust. *International Journal of Social Robotics*, 13(8), 2021–2038, [žiūrēta 2024-11-15]. Prieiga per internetą: <https://doi.org/10.1007/s12369-020-00692-3>
 17. Banks, J., Koban, K., & Haggadone, B. (2023). Avoiding the Abject and Seeking the Script: Perceived Mind, Morality, and Trust in a Persuasive Social Robot. *ACM Transactions on Human-Robotic Interaction*, 12(3), 1–24, [žiūrēta 2024-12-10]. Prieiga per internetą: <https://doi.org/10.1145/3572036>
 18. Bigman, Y. E., Wilson, D., Arnestad, M. N., Waytz, A., & Gray, K. (2022). Algorithmic discrimination causes less moral outrage than human discrimination. *Journal of Experimental Psychology*, 151(12), 2683–2700, [žiūrēta 2025-02-08]. Prieiga per internetą: <https://doi.org/10.1037/xge0001250>
 19. Brailsford, J., Vetere, F., & Velloso, E. (2024). Exploring the association between moral foundations and judgements of AI behaviour. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, 284, 1–15, [žiūrēta 2025-02-22]. Prieiga per internetą: <https://doi.org/10.1145/3613904.3642712>
 20. Brandtner, P., Darbanian, F., Falatouri, T., & Udokwu, C. (2021). Impact of COVID-19 on the Customer End of Retail Supply Chains: A Big Data Analysis of Consumer Satisfaction. *Sustainability*, 13(3), 1464, [žiūrēta 2025-03-04]. Prieiga per internetą: <https://doi.org/10.3390/su13031464>
 21. Brendel, A. B., Mirbabaie, M., Lembcke, T.-B., & Hofeditz, L. (2021). Ethical Management of Artificial Intelligence. *Sustainability*, 13(4), 1974, [žiūrēta 2025-01-21]. Prieiga per internetą: <https://doi.org/10.3390/su13041974>
 22. Choi, J., & Chao, M. M. (2024). For Me or Against Me? Reactions to AI (vs. Human) Decisions That Are Favorable or Unfavorable to the Self and the Role of Fairness Perception. *Personality & Social Psychology Bulletin*, 1-21, [žiūrēta 2025-02-23]. Prieiga per internetą: <https://doi.org/10.1177/01461672241288338>
 23. Choung, H., Seberger, J. S., & David, P. (2023). When AI is Perceived to Be Fairer than a Human: Understanding Perceptions of Algorithmic Decisions in a Job Application Context. *International Journal of Human-Computer Interaction*, 40(22), 7451–7468, [žiūrēta 2025-02-24]. Prieiga per internetą: <https://doi.org/10.1080/10447318.2023.2266244>

24. Christ-Brendemühl, S., & Schaarschmidt, M. (2022). Customer fairness perceptions in augmented reality-based online services. *Journal of Service Management*, 33(1), 9–32, [žiūrėta 2025-03-09]. Prieiga per internetą: <https://doi.org/10.1108/JOSM-01-2021-0012>
25. Clegg, M., Hofstetter, R., de Bellis, E., & Schmitt, B. H. (2024). Unveiling the Mind of the Machine. *The Journal of Consumer Research*, 51(2), 342–361, [žiūrėta 2025-02-12]. Prieiga per internetą: <https://doi.org/10.1093/jcr/ucad075>
26. Coeckelbergh, M. (2016). Responsibility and the Moral Phenomenology of Using Self-Driving Cars. *Applied Artificial Intelligence*, 30(8), 748–757, [žiūrėta 2024-12-19]. Prieiga per internetą: <https://doi.org/10.1080/08839514.2016.1229759>
27. Constantinescu, M., Vică, C., Uszkai, R., & Voinea, C. (2022). Blame It on the AI? On the Moral Responsibility of Artificial Moral Advisors. *Philosophy & Technology*, 35(2), [žiūrėta 2025-01-03]. Prieiga per internetą: <https://doi.org/10.1007/s13347-022-00529-z>
28. De Brigard, F. (2023). Counterfactual Thinking. In: Glăveanu, V.P. (eds) *The Palgrave Encyclopedia of the Possible*. Palgrave Macmillan, 243–250, [žiūrėta 2025-02-25]. Prieiga per internetą: https://doi.org/10.1007/978-3-030-90913-0_43
29. de Pedro, J. M. L. (2024). Social Systems as Moral Agents: A Systems Approach to Moral Agency in Business. *Journal of Business Ethics*, 195(4), 695–711, [žiūrėta 2024-12-01]. Prieiga per internetą: <https://doi.org/10.1007/s10551-024-05677-0>
30. Di Dio, C., Manzi, F., Peretti, G., Cangelosi, A., Harris, P. L., Massaro, D., & Marchetti, A. (2020). Shall I Trust You? From Child-Robot Interaction to Trusting Relationships. *Frontiers in psychology*, 11, 469, [žiūrėta 2025-02-18]. Prieiga per internetą: <https://doi.org/10.3389/fpsyg.2020.00469>
31. Dietvorst, B. J., & Bartels, D. M. (2022). Consumers Object to Algorithms Making Morally Relevant Tradeoffs Because of Algorithms' Consequentialist Decision Strategies. *Journal of Consumer Psychology*, 32(3), 406–424, [žiūrėta 2025-02-09]. Prieiga per internetą: <https://doi.org/10.1002/jcpy.1266>
32. Dong, M., & Bocian, K. (2024). Responsibility gaps and self-interest bias: People attribute moral responsibility to AI for their own but not others' transgressions. *Journal of Experimental Social Psychology*, 111, 104584, [žiūrėta 2024-12-19]. Prieiga per internetą: <https://doi.org/10.1016/j.jesp.2023.104584>
33. Douglas, B. D., Ewell, P. J., & Brauer, M. (2023). Data quality in online human-subjects research: Comparisons between MTurk, Prolific, CloudResearch, Qualtrics, and SONA. *PloS One*, 18(3), e0279720–e0279720, [žiūrėta 2025-03-23]. Prieiga per internetą: <https://doi.org/10.1371/journal.pone.0279720>
34. Du, S., & Xie, C. (2021). Paradoxes of artificial intelligence in consumer markets: Ethical challenges and opportunities. *Journal of Business Research*, 129, 961–974, [žiūrėta 2025-01-19]. Prieiga per internetą: <https://doi.org/10.1016/j.jbusres.2020.08.024>
35. Ebrahimi, S., Abdelhalim, E., Hassanein, K., & Head, M. (2024). Reducing the incidence of biased algorithmic decisions through feature importance transparency: an empirical study. *European Journal of Information Systems*, 1–29, [žiūrėta 2025-03-03]. Prieiga per internetą: <https://doi.org/10.1080/0960085X.2024.2395531>
36. Esterwood, C., & Robert, L. P. (2023). The theory of mind and human–robot trust repair. *Scientific Reports*, 13(1), 9877–15, [žiūrėta 2025-02-10]. Prieiga per internetą: <https://doi.org/10.1038/s41598-023-37032-0>

37. Firt, E. (2024). What makes full artificial agents morally different. *AI & Society*, [žiūrėta 2024-12-18]. Prieiga per internetą: <https://doi.org/10.1007/s00146-024-01867-6>
38. Formosa, P., & Ryan, M. (2021). Making moral machines: why we need artificial moral agents. *AI & Society*, 36(3), 839–851, [žiūrėta 2024-12-18]. Prieiga per internetą: <https://doi.org/10.1007/s00146-020-01089-6>
39. Galvez, V., & Hanono, E. (2024). What Does it Mean to Measure Mind Perception toward Robots? A Critical Review of the Main Self-Report Instruments. *International Journal of Social Robotics*, 16(3), 501–511, [žiūrėta 2025-02-23]. Prieiga per internetą: <https://doi.org/10.1007/s12369-024-01113-5>
40. Gamez, P., Shank, D. B., Arnold, C., & North, M. (2020). Artificial virtue: the machine question and perceptions of moral character in artificial moral agents. *AI & Society*, 35(4), 795–809, [žiūrėta 2024-12-19]. Prieiga per internetą: <https://doi.org/10.1007/s00146-020-00977-1>
41. Garvey, A. M., Kim, T., & Duhachek, A. (2023). Bad News? Send an AI. Good News? Send a Human. *Journal of Marketing*, 87(1), 10–25, [žiūrėta 2024-11-15]. Prieiga per internetą: <https://doi.org/10.1177/00222429211066972>
42. Geiselmann, R., Tsourgianni, A., Deroy, O., & Harris, L. T. (2023). Interacting with agents without a mind: the case for artificial agents. *Current Opinion in Behavioral Sciences*, 51, 101282, [žiūrėta 2025-02-09]. Prieiga per internetą: <https://doi.org/10.1016/j.cobeha.2023.101282>
43. Gill, T. (2020). Blame It on the Self-Driving Car: How Autonomous Vehicles Can Alter Consumer Morality. *The Journal of Consumer Research*, 47(2), 272–291, [žiūrėta 2025-02-09]. Prieiga per internetą: <https://doi.org/10.1093/jcr/ucaa018>
44. Goenka, S., & Thomas, M. (2024). Moral foundations theory and consumer behavior. *Journal of Consumer Psychology*, 34(3), 536–540, [žiūrėta 2025-02-22]. Prieiga per internetą: <https://doi.org/10.1002/jcpy.1429>
45. Graham, J., Nosek, B., Haidt, J., Iyer, R., Koleva, S., & Ditto, P. (2011). Mapping the moral domain. *Journal of Personality and Social Psychology*, 101(2), 366–385, [žiūrėta 2025-02-22]. Prieiga per internetą: <https://doi.org/10.1037/a0021847>.
46. Gray, H. M., Gray, K., & Wegner, D. M. (2007). Dimensions of Mind Perception. *Science (American Association for the Advancement of Science)*, 315(5812), 619–619, [žiūrėta 2025-02-17]. Prieiga per internetą: <https://doi.org/10.1126/science.1134475>
47. Grgić-Hlača, N., Lima, G., Weller, A., & Redmiles, E. M. (2022). Dimensions of diversity in human perceptions of algorithmic fairness. In *Proceedings of the 2nd ACM Conference on Equity and Access in Algorithms, Mechanisms, and Optimization (EAAMO '22)*, 21, 1–12, [žiūrėta 2025-03-02]. Prieiga per internetą: <https://doi.org/10.1145/3551624.3555306>
48. Gummer, T., Roßmann, J., & Silber, H. (2021). Using Instructed Response Items as Attention Checks in Web Surveys: Properties and Implementation. *Sociological Methods & Research*, 50(1), 238–264, [žiūrėta 2025-04-07]. Prieiga per internetą: <https://doi.org/10.1177/0049124118769083>
49. Gupta, M., Parra, C. M., & Dennehy, D. (2022). Questioning Racial and Gender Bias in AI-based Recommendations: Do Espoused National Cultural Values Matter? *Information Systems Frontiers*, 24(5), 1465–1481, [žiūrėta 2025-03-02]. Prieiga per internetą: <https://doi.org/10.1007/s10796-021-10156-2>
50. Habes, M., Ali, S., & Pasha, S. A. (2021). Statistical Package for Social Sciences Acceptance in Quantitative Research: From the Technology Acceptance Model's Perspective. (2021). *FWU*

- Journal of Social Sciences*, 34–46, [žiūrēta 2024-11-19]. Prieiga per internetą: <https://doi.org/10.51709/19951272/Winter-2021/3>
51. Haidt, J., Ditto, P., Iyer, R., Wojcik, S., Koleva, S., Motyl, M., & Graham, J. (2013). Chapter Two - Moral Foundations Theory: The Pragmatic Validity of Moral Pluralism. In *Advances in Experimental Social Psychology* (Vol. 47, pp. 55–130). Elsevier Science & Technology. <https://www.elsevier.com/>
 52. Harris, J., & Anthis, J. R. (2021). The Moral Consideration of Artificial Entities: A Literature Review. *Science and Engineering Ethics*, 27(4), 53, [žiūrēta 2024-12-18]. Prieiga per internetą: <https://doi.org/10.1007/s11948-021-00331-8>
 53. Haupt, M., Freidank, J., & Haas, A. (2024). Consumer responses to human-AI collaboration at organizational frontlines: strategies to escape algorithm aversion in content creation. *Review of Managerial Science*. [žiūrēta 2025-01-22]. Prieiga per internetą: <https://doi.org/10.1007/s11846-024-00748-y>
 54. Helberger, N., Araujo, T., & de Vreese, C. H. (2020). Who is the fairest of them all? Public attitudes and expectations regarding automated decision-making. *The Computer Law and Security Report*, 39, 105456, [žiūrēta 2025-02-24]. Prieiga per internetą: <https://doi.org/10.1016/j.clsr.2020.105456>
 55. Huang, M.-H., & Rust, R. T. (2021). Engaged to a Robot? The Role of AI in Service. *Journal of Service Research: JSR*, 24(1), 30–41, [žiūrēta 2024-12-07]. Prieiga per internetą: <https://doi.org/10.1177/1094670520902266>
 56. Hwang, A. H.-C., & Won, A. S. (2022). AI in your mind: Counterbalancing perceived agency and experience in human-AI interaction. *Extended Abstracts of the 2022 CHI Conference on Human Factors in Computing Systems (CHI EA '22)*, 1–10. Association for Computing Machinery. <https://doi.org/10.1145/3491101.3519833>
 57. Yalcin, G., Lim, S., Puntoni, S., & van Osselaer, S. M. J. (2022). Thumbs Up or Down: Consumer Reactions to Decisions by Algorithms Versus Humans. *Journal of Marketing Research*, 59(4), 696–717, [žiūrēta 2025-01-24]. Prieiga per internetą: <https://doi.org/10.1177/00222437211070016>
 58. Yam, K.C., Bigman, Y.E., Tang, P.M., Ilies, R., De Cremer, D., Soh, H., & Gray, K. (2021). Robots at work: People prefer-and forgive-service robots with perceived feelings. *Journal of applied psychology*, 106(10), 1557–1572, [žiūrēta 2025-02-12]. Prieiga per internetą: <https://doi.org/10.1037/apl0000834>
 59. Yang, Q., & Lee, Y.-C. (2024). Ethical AI in Financial Inclusion: The Role of Algorithmic Fairness on User Satisfaction and Recommendation. *Big Data and Cognitive Computing*, 8(9), 105-, [žiūrēta 2025-03-10]. Prieiga per internetą: <https://doi.org/10.3390/bdcc8090105>
 60. Yoo, K., Haenlein, M., & Hewett, K. (2025). A whole new world, a new fantastic point of view: Charting unexplored territories in consumer research with generative artificial intelligence. *Journal of the Academy of Marketing Science*, [žiūrēta 2025-04-10]. Prieiga per internetą: <https://doi.org/10.1007/s11747-025-01097-2>
 61. Jain, V., Wadhwani, K., & Eastman, J. K. (2024). Artificial intelligence consumer behavior: A hybrid review and research agenda. *Journal of Consumer Behaviour*, 23(2), 676–697, [žiūrēta 2025-02-23]. Prieiga per internetą: <https://doi.org/10.1002/cb.2233>

62. Jimenez Mori, R. (2021). It's not price; It's quality. Satisfaction and price fairness perception. *World Development*, 139, 105302, [žiūrėta 2025-03-08]. Prieiga per internetą: <https://doi.org/10.1016/j.worlddev.2020.105302>
63. Jung, H., Bae, J., & Kim, H. (2022). The effect of corporate social responsibility and corporate social irresponsibility: Why company size matters based on consumers' need for self-expression. *Journal of Business Research*, 146, 146–154, [žiūrėta 2025-02-25]. Prieiga per internetą: <https://doi.org/10.1016/j.jbusres.2022.03.024>
64. Koh, L. Y., Cheh, Y. G., Yuen, K. F., & Wang, X. (2025). Satisfaction with AI chatbots in crowdsourcing services: anthropomorphic and fairness perspectives. *Behaviour & Information Technology*, 1–22, [žiūrėta 2025-03-08]. Prieiga per internetą: <https://doi.org/10.1080/0144929X.2025.2469660>
65. Laakasuo, M., Palomäki, J., & Köbis, N. (2021). Moral Uncanny Valley: A Robot's Appearance Moderates How its Decisions are Judged. *International Journal of Social Robotics*, 13(7), 1679–1688, [žiūrėta 2024-11-14]. Prieiga per internetą: <https://doi.org/10.1007/s12369-020-00738-6>
66. Ladak, A., Loughnan, S., & Wilks, M. (2024). The Moral Psychology of Artificial Intelligence. *Current Directions in Psychological Science : A Journal of the American Psychological Society*, 33(1), 27–34, [žiūrėta 2025-02-10]. Prieiga per internetą: <https://doi.org/10.1177/09637214231205866>
67. Leclercq-Machado, L., Alvarez-Risco, A., Esquerre-Botton, S., Almanza-Cruz, C., de las Mercedes Anderson-Seminario, M., Del-Aguila-Arcntales, S., & Yáñez, J. A. (2022). Effect of Corporate Social Responsibility on Consumer Satisfaction and Consumer Loyalty of Private Banking Companies in Peru. *Sustainability*, 14(15), 9078, [žiūrėta 2025-03-04]. Prieiga per internetą: <https://doi.org/10.3390/su14159078>
68. Lee, I., & Hahn, S. (2024). On the relationship between mind perception and social support of chatbots. *Frontiers in Psychology*, 15, 1282036, [žiūrėta 2025-02-12]. Prieiga per internetą: <https://doi.org/10.3389/fpsyg.2024.1282036>
69. Lee, M. K. (2018). Understanding perception of algorithmic decisions: Fairness, trust, and emotion in response to algorithmic management. *Big Data & Society*, 5(1), [žiūrėta 2025-03-02]. Prieiga per internetą: <https://doi.org/10.1177/2053951718756684>
70. Lee, M., Ruijten, P., Frank, L., de Kort, Y., & IJsselstein, W. (2021). People may punish, but not blame robots. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems (CHI '21)*, 715, [žiūrėta 2025-02-17]. Prieiga per internetą: <https://doi.org/10.1145/3411764.3445284>
71. Li, S., Peluso, A. M., & Duan, J. (2023). Why do we prefer humans to artificial intelligence in telemarketing? A mind perception explanation. *Journal of Retailing and Consumer Services*, 70, 103139, [žiūrėta 2025-02-12]. Prieiga per internetą: <https://doi.org/10.1016/j.jretconser.2022.103139>
72. Li, Z., Terfurth, L., Woller, J. P., & Wiese, E. (2022). Mind the machines: Applying implicit measures of mind perception to social robotics. *Proceedings of the 17th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*, 236–245, [žiūrėta 2025-02-12]. Prieiga per internetą: <https://10.1109/HRI53351.2022.9889356>
73. Liu, Y., & Sun, X. (2024). Towards more legitimate algorithms: A model of algorithmic ethical perception, legitimacy, and continuous usage intentions of e-commerce platforms. *Computers in*

- Human Behavior*, 150, 108006-, [žiūrėta 2025-03-10]. Prieiga per internetą: <https://doi.org/10.1016/j.chb.2023.108006>
74. Liu, J. T., & Mendez, M. J. (2022). The Attributional-Counterfactual Theory of Need: Integrating Theories to Predict Need Norm Use. *Business Ethics Quarterly*, 32(1), 103–135, [žiūrėta 2025-02-25]. Prieiga per internetą: <https://doi.org/10.1017/beq.2021.1>
 75. Llorca Albareda, J. (2024). Anthropological Crisis or Crisis in Moral Status: a Philosophy of Technology Approach to the Moral Consideration of Artificial Intelligence. *Philosophy & Technology*, 37(1), 12, [žiūrėta 2024-12-14]. Prieiga per internetą: <https://doi.org/10.1007/s13347-023-00682-z>
 76. Longoni, C., Bonezzi, A., & Morewedge, C. (2019). Resistance to Medical Artificial Intelligence. *Journal of Consumer Research*, 46, 629–650. [žiūrėta 2025-01-22]. [Phttps://doi.org/10.1093/jcr/ucz013](https://doi.org/10.1093/jcr/ucz013)
 77. Luttrell, A., Teeny, J. D., & Petty, R. E. (2021). Morality matters in the marketplace: The role of moral metacognition in consumer purchasing. *Social Cognition*, 39 (3), 328–351 [žiūrėta 2025-02-22]. Prieiga per internetą: <https://doi.org/10.1521/soco.2021.39.3.328>
 78. Maninger, T., & Shank, D. B. (2022). Perceptions of violations by artificial and human actors across moral foundations. *Computers in Human Behavior Reports*, 5, 100154, [žiūrėta 2024-12-19]. Prieiga per internetą: <https://doi.org/10.1016/j.chbr.2021.100154>
 79. Manoli, A., Pauketat, J. V. T., & Anthis, J. R. (2024). The AI double standard: Humans judge all AIs for the actions of one. *Proceedings of the ACM on Human-Computer Interaction*, 9(CSCW), [žiūrėta 2025-02-01]. Prieiga per internetą: <https://doi.org/10.48550/arXiv.2412.06040>
 80. Mariani, M. M., Perez-Vega, R., & Wirtz, J. (2022). AI in marketing, consumer research and psychology: A systematic literature review and research agenda. *Psychology & Marketing*, 39(4), 755–776, [žiūrėta 2025-02-17]. Prieiga per internetą: <https://doi.org/10.1002/mar.21619>
 81. Mathur, P., & Sarin Jain, S. (2020). Not all that glitters is golden: The impact of procedural fairness perceptions on firm evaluations and customer satisfaction with favorable outcomes. *Journal of Business Research*, 117, 357–367, [žiūrėta 2025-02-26]. Prieiga per internetą: <https://doi.org/10.1016/j.jbusres.2020.06.006>
 82. Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., & Galstyan, A. (2021). A Survey on Bias and Fairness in Machine Learning. *ACM Computing Surveys*, 54(6), 1–35, [žiūrėta 2025-02-24]. Prieiga per internetą: <https://doi.org/10.1145/3457607>
 83. Mohajan, H. K. (2020). Quantitative Research: A Successful Investigation in Natural and Social Sciences. *Journal of Economic Development, Environment and People*, 9(4), 50–79, [žiūrėta 2025-03-24]. Prieiga per internetą: <https://doi.org/10.26458/jedep.v9i4.679>
 84. Müller, B. C. N., Gao, X., Nijssen, S. R. R., & Damen, T. G. E. (2021). I, Robot: How Human Appearance and Mind Attribution Relate to the Perceived Danger of Robots. *International Journal of Social Robotics*, 13(4), 691–701, [žiūrėta 2025-02-12]. Prieiga per internetą: <https://doi.org/10.1007/s12369-020-00663-8>
 85. Narayanan, D., Nagpal, M., McGuire, J., Schweitzer, S., & De Cremer, D. (2024). Fairness Perceptions of Artificial Intelligence: A Review and Path Forward. *International Journal of Human-Computer Interaction*, 40(1), 4–23, [žiūrėta 2025-01-25]. Prieiga per internetą: <https://doi.org/10.1080/10447318.2023.2210890>

86. Nicolitsas, D. (2024). Fairness, expectations and life satisfaction: evidence from Europe. *Empirica*, 51(2), 313–349. [žiūrėta 2025-03-08]. Prieiga per internetą: <https://doi.org/10.1007/s10663-023-09602-y>
87. Nijssen, S. R. R., Müller, B. C. N., Bosse, T., & Paulus, M. (2023). Can you count on a calculator? The role of agency and affect in judgments of robots as moral agents. *Human-Computer Interaction*, 38(5–6), 400–416, [žiūrėta 2025-02-03]. Prieiga per internetą: <https://doi.org/10.1080/07370024.2022.2080552>
88. Noble, S. M., Foster, L. L., & Craig, S. B. (2021). The procedural and interpersonal justice of automated application and resume screening. *International Journal of Selection and Assessment*, 29(2), 139–153, [žiūrėta 2025-03-03]. Prieiga per internetą: <https://doi.org/10.1111/ijsa.12320>
89. Obermeyer, Z., Powers, B., Vogeli, C., & Mullainathan, S. (2019). Dissecting racial bias in an algorithm used to manage the health of populations. *Science (American Association for the Advancement of Science)*, 366(6464), 447–453, [žiūrėta 2025-04-10]. Prieiga per internetą: <https://doi.org/10.1126/science.aax2342>
90. Ochmann, J., Michels, L., Tiefenbeck, V., Maier, C., & Laumer, S. (2024). Perceived algorithmic fairness: An empirical study of transparency and anthropomorphism in algorithmic recruiting. *Information Systems Journal (Oxford, England)*, 34(2), 384–414, [žiūrėta 2025-02-27]. Prieiga per internetą: <https://doi.org/10.1111/isj.12482>
91. Pallant, J. (2010). *SPSS survival manual : a step by step guide to data analysis using SPSS* / Julie Pallant. (4th ed.). McGraw Hill.
92. Pavone, G., Meyer-Waarden, L., & Munzel, A. (2023). Rage Against the Machine: Experimental Insights into Customers' Negative Emotional Responses, Attributions of Responsibility, and Coping Strategies in Artificial Intelligence–Based Service Failures. *Journal of Interactive Marketing*, 58(1), 52–71, [žiūrėta 2024-11-15]. Prieiga per internetą: <https://doi.org/10.1177/10949968221134492>
93. Pitardi, V., Wirtz, J., Paluch, S., & Kunz, W. H. (2022). Service robots, agency and embarrassing service encounters. *Journal of Service Management*, 33(2), 389–414, [žiūrėta 2024-11-13]. Prieiga per internetą: <https://doi.org/10.1108/JOSM-12-2020-0435>
94. Qin, S. (Marco), Jia, N., Luo, X., Liao, C., & Huang, Z. (2023). Perceived Fairness of Human Managers Compared with Artificial Intelligence in Employee Performance Evaluation. *Journal of Management Information Systems*, 40(4), 1039–1070. [žiūrėta 2025-03-03]. Prieiga per internetą: <https://doi.org/10.1080/07421222.2023.2267316>
95. Ramos, G. A., Johnson, W., VanEpps, E. M., & Graham, J. (2024). When consumer decisions are moral decisions: Moral Foundations Theory and its implications for consumer psychology. *Journal of Consumer Psychology*, 34(3), 519–535, [žiūrėta 2025-02-22]. Prieiga per internetą: <https://doi.org/10.1002/jcpy.1427>
96. Ryoo, Y., Jeon, Y. A., & Kim, W. (2024). The blame shift: Robot service failures hold service firms more accountable. *Journal of Business Research*, 171, 114360, [žiūrėta 2024-11-11]. Prieiga per internetą: <https://doi.org/10.1016/j.jbusres.2023.114360>
97. Sarstedt, M., Adler, S. J., Rau, L., & Schmitt, B. (2024). *Psychology & Marketing*, 41(6), 1254–1270, [žiūrėta 2025-04-10]. Prieiga per internetą: <https://doi.org/10.1002/mar.21982>
98. Schoeffer, J., Machowski, Y., & Kuehl, N. (2021a). A study on fairness and trust perceptions in automated decision making. In *Joint Proceedings of the ACM IUI 2021 Workshops*, 1–12, [žiūrėta 2025-03-02]. Prieiga per internetą: <https://doi.org/10.48550/arXiv.2103.04757>

99. Schoeffer, J., Machowski, Y., & Kuehl, N. (2021b). Perceptions of fairness and trustworthiness based on explanations in human vs. automated decision-making. *Proceedings of the 55th Hawaii International Conference on System Sciences*, 1–8, [žiūrėta 2025-03-03]. Prieiga per internetą: <http://hdl.handle.net/10125/79466>
100. Scholl-Grissemann, U., Stokburger-Sauer, N. E., & Teichmann, K. (2022). The importance of perceived fairness in product customization settings. *The Service Industries Journal*, 42(11–12), 823–842, [žiūrėta 2025-03-02]. Prieiga per internetą: <https://doi.org/10.1080/02642069.2020.1819252>
101. Setiawan, E. B., Wati, S., Wardana, A., & Ikhsan, R. B. (2020). Building trust through customer satisfaction in the airline industry in Indonesia: Service quality and price fairness contribution. *Management Science Letters*, 10(5), 1095–1102, [žiūrėta 2025-03-09]. Prieiga per internetą: <https://doi.org/10.5267/j.msl.2019.10.033>
102. Shank, D. B., DeSanti, A., & Maninger, T. (2019). When are artificial intelligence versus human agents faulted for wrongdoing? Moral attributions after individual and joint decisions. *Information, Communication & Society*, 22(5), 648–663, [žiūrėta 2025-02-08]. Prieiga per internetą: <https://doi.org/10.1080/1369118X.2019.1568515>
103. Shank, D., North, M., Arnold, C., & Gamez, P. (2021). Can mind perception explain virtuous character judgments of artificial intelligence? *Technology, Mind, and Behavior*, 2(1), 1–13 [žiūrėta 2025-02-12]. Prieiga per internetą: <https://doi.org/10.1037/tmb0000047.supp>
104. Shin, D. (2021). The effects of explainability and causability on perception, trust, and acceptance: Implications for explainable AI. *International Journal of Human-Computer Studies*, 146, 102551–, [žiūrėta 2025-03-03]. Prieiga per internetą: <https://doi.org/10.1016/j.ijhcs.2020.102551>
105. Shulner-Tal, A., Kuflik, T., & Kliger, D. (2023). Enhancing Fairness Perception - Towards Human-Centred AI and Personalized Explanations Understanding the Factors Influencing Laypeople's Fairness Perceptions of Algorithmic Decisions. *International Journal of Human-Computer Interaction*, 39(7), 1455–1482, [žiūrėta 2025-01-24]. Prieiga per internetą: <https://doi.org/10.1080/10447318.2022.2095705>
106. Sison, A. J. G., & Redín, D. M. (2023). A neo-aristotelian perspective on the need for artificial moral agents (AMAs). *AI & Society*, 38(1), 47–65, [žiūrėta 2024-12-20]. Prieiga per internetą: <https://doi.org/10.1007/s00146-021-01283-0>
107. Söderlund, M. (2022). Service robots with (perceived) theory of mind: An examination of humans' reactions. *Journal of Retailing and Consumer Services*, 67, 102999, [žiūrėta 2025-02-12]. Prieiga per internetą: <https://doi.org/10.1016/j.jretconser.2022.102999>
108. Söderlund, M. (2023a). Service robot verbalization in service processes with moral implications and its impact on satisfaction. *Technological Forecasting & Social Change*, 196, 122831–, [žiūrėta 2024-12-21]. Prieiga per internetą: <https://doi.org/10.1016/j.techfore.2023.122831>
109. Soderlund, M. (2023b). Service robots and artificial morality: an examination of robot behavior that violates human privacy. *Journal of Service Theory and Practice*, 33(7), 52–72, [žiūrėta 2025-02-09]. Prieiga per internetą: <https://doi.org/10.1108/JSTP-09-2022-0196>
110. Söderlund, M. (2024). Human employees and service robots in the service encounter and the role of attribution of theory of mind. *Journal of Retailing and Consumer Services*, 81, 103964, [žiūrėta 2025-02-17]. Prieiga per internetą: <https://doi.org/10.1016/j.jretconser.2024.103964>

111. Starke, C., Baleis, J., Keller, B., & Marcinkowski, F. (2022). Fairness perceptions of algorithmic decision-making: A systematic review of the empirical literature. *Big Data & Society*, 9(2), [žiūrėta 2025-03-01]. Prieiga per internetą: <https://doi.org/10.1177/20539517221115189>
112. Stein, J.P., Appel, M., Jost, A., & Ohler, P. (2020). Matter over mind? How the acceptance of digital entities depends on their appearance, mental prowess, and the interaction between both. *International Journal of Human-Computer Studies*, 142, 102463, [žiūrėta 2025-02-12]. Prieiga per internetą: <https://doi.org/10.1016/j.ijhcs.2020.102463>
113. Stoner, J. L., Felix, R., & Stadler Blank, A. (2023). Best practices for implementing experimental research methods. *International Journal of Consumer Studies*, 47(4), 1579–1595, [žiūrėta 2025-03-23]. Prieiga per internetą: <https://doi.org/10.1111/ijcs.12878>
114. Sullivan, Y. W., & Fosso Wamba, S. (2022). Moral Judgments in the Age of Artificial Intelligence. *Journal of Business Ethics*, 178(4), 917–943, [žiūrėta 2024-11-09]. Prieiga per internetą: <https://doi.org/10.1007/s10551-022-05053-w>
115. Sürücü, L., Şeşen, H., & Maslakçı, A. (2023). Regression, mediation/moderation, and structural equation modeling with SPSS, AMOS, and PROCESS Macro. *Livre de Lyon*.
116. Swiderska, A., & Küster, D. (2020). Robots as Malevolent Moral Agents: Harmful Behavior Results in Dehumanization, Not Anthropomorphism. *Cognitive Science*, 44(7), e12872-n/a, [žiūrėta 2025-02-09]. Prieiga per internetą: <https://doi.org/10.1111/cogs.12872>
117. Telkamp, J. B., & Anderson, M. H. (2022). The Implications of Diverse Human Moral Foundations for Assessing the Ethicality of Artificial Intelligence. *Journal of Business Ethics*, 178(4), 961–976, [žiūrėta 2025-02-22]. Prieiga per internetą: <https://doi.org/10.1007/s10551-022-05057-6>
118. Thellman, S., de Graaf, M., & Ziemke, T. (2022). Mental state attribution to robots: A systematic review of conceptions, methods, and findings. *Journal of Human-Robot Interaction*, 11(4), 41, [žiūrėta 2025-02-11]. Prieiga per internetą: <https://doi.org/10.1145/3526112>
119. Trafton, J. G., Frazier, C. R., Zish, K., Bio, B. J., & McCurry, J. M. (2023). The Perception of Agency: Scale Reduction and Construct Validity. In *Proceedings of the 2023 32nd IEEE International Conference on Robot and Human Interactive Communication (RO-MAN)*, 936–942, [žiūrėta 2025-02-03]. Prieiga per internetą: <https://doi.org/10.1109/RO-MAN57019.2023.10309544>
120. van Berkel, N., Goncalves, J., Russo, D., Hosio, S., & Skov, M. B. (2021). Effect of information presentation on fairness perceptions of machine learning predictors. In *CHI Conference on Human Factors in Computing Systems*, 245, 1–13, [žiūrėta 2025-03-02]. Prieiga per internetą: <https://doi.org/10.1145/3411764.344536>
121. Van Quaquebeke, N., Salem, M., van Dijke, M., & Wenzel, R. (2022). Conducting organizational survey and experimental research online: From convenient to ambitious in study designs, recruiting, and data quality. *Organizational Psychology Review*, 12(3), 268–305, [žiūrėta 2025-03-24]. Prieiga per internetą: <https://doi.org/10.1177/20413866221097571>
122. Véliz, C. (2021). Moral zombies: why algorithms are not moral agents. *AI & Society*, 36(2), 487–497, [žiūrėta 2024-12-19]. Prieiga per internetą: <https://doi.org/10.1007/s00146-021-01189-x>
123. Viglia, G., Zaefarian, G., & Ulqinaku, A. (2021). How to design good experiments in marketing: Types, examples, and methods. *Industrial Marketing Management*, 98, 193–206. [žiūrėta 2025-04-07]. Prieiga per internetą: <https://doi.org/10.1016/j.indmarman.2021.08.007>

124. Vijayaraghavan, A., & Badea, C. (2024). Minimum levels of interpretability for artificial moral agents. *Ai and Ethics (Online)*, [žiūrēta 2024-12-07]. Prieiga per internetą: <https://doi.org/10.1007/s43681-024-00536-0>
125. Wang, X., Yuen, K. F., Teo, C.-C., & Wong, Y. D. (2021). Online Consumers' Satisfaction in Self-Collection: Value Co-Creation from the Service Fairness Perspective. *International Journal of Electronic Commerce*, 25(2), 230–260, [žiūrēta 2025-03-08]. Prieiga per internetą: <https://doi.org/10.1080/10864415.2021.1887699>
126. Wang, X., Zhang, Y., & Zhu, R. (2022). A brief review on algorithmic fairness. *Management System Engineering*, 1(1), [žiūrēta 2025-02-24]. Prieiga per internetą: <https://doi.org/10.1007/s44176-022-00006-z>
127. Weith, H., & Matt, C. (2022). When do customers perceive artificial intelligence as fair? An assessment of AI-based B2C e-commerce. *Proceedings of the 55th Hawaii International Conference on System Sciences*, 4336–4345, [žiūrēta 2025-02-28]. Prieiga per internetą: I: <https://hdl.handle.net/10125/79866>
128. Wester, J., Pohl, H., Hosio, S., & van Berkel, N. (2024). “This chatbot would never...”: Perceived moral agency of mental health chatbots. *Proceedings of the ACM on Human-Computer Interaction*, 8(CSCW1), 133, [žiūrēta 2025-01-28]. Prieiga per internetą: <https://doi.org/10.1145/3637410>
129. Wilson, A., Stefanik, C., & Shank, D. B. (2022). How do people judge the immorality of artificial intelligence versus humans committing moral wrongs in real-world situations? *Computers in Human Behavior Reports*, 8, 100229, [žiūrēta 2025-02-08]. Prieiga per internetą: <https://doi.org/10.1016/j.chbr.2022.100229>
130. Wu, X., Li, S., Guo, Y., & Fang, S. (2024). Human or AI robot? Who is fairer on the service organizational frontline. *Journal of Business Research*, 181, 114730, [žiūrēta 2025-01-26]. Prieiga per internetą: <https://doi.org/10.1016/j.jbusres.2024.114730>
131. Xu, Y., Shieh, C.-H., van Esch, P., & Ling, I.-L. (2020). AI customer service: Task complexity, problem-solving ability, and usage intention. *Australasian Marketing Journal*, 28(4), 189–199, [žiūrēta 2024-12-17]. Prieiga per internetą: <https://doi.org/10.1016/j.ausmj.2020.03.005>
132. Zafar, M. (2024). Normativity and AI moral agency. *Ai and Ethics (Online)*, [žiūrēta 2024-12-14]. Prieiga per internetą: <https://doi.org/10.1007/s43681-024-00566-8>
133. Zafari, S., & Koeszegi, S. T. (2021). Attitudes Toward Attributed Agency: Role of Perceived Control. *International Journal of Social Robotics*, 13(8), 2071–2080, [žiūrēta 2025-02-09]. Prieiga per internetą: <https://doi.org/10.1007/s12369-020-00672-7>
134. Zakharin, M., & Bates, T. C. (2023). Moral Foundations Theory: Validation and replication of the MFQ-2. *Personality and Individual Differences*, 214, 112339, [žiūrēta 2025-02-22]. Prieiga per internetą: <https://doi.org/10.1016/j.paid.2023.112339>
135. Zhang, Y., Wu, J., Yu, F., & Xu, L. (2023). Moral Judgments of Human vs. AI Agents in Moral Dilemmas. *Behavioral Sciences*, 13(2), 181, [žiūrēta 2024-12-19]. Prieiga per internetą: <https://doi.org/10.3390/bs13020181>
136. Zhao, Y., Zhu, Z., & Tang, B. (2025). Are highly anthropomorphic service robots more likely to be forgiven by customers after service failures? A mind perception perspective. *International Journal of Hospitality Management*, 126, 104103, [žiūrēta 2025-02-18]. Prieiga per internetą: <https://doi.org/10.1016/j.ijhm.2025.104103>

Informacijos šaltinių sąrašas

1. BCG. (2025). AI radar 2025. Boston Consulting Group. [žiūrėta 2025-01-21]. Prieiga per internetą: <https://web-assets.bcg.com/0b/f6/c2880f9f4472955538567a5bcb6a/ai-radar-2025-slideshow-jan-2025-r.pdf>
2. Global Market Insights. (2024). Salon Service Market Size - By Service (Hair Care, Nail Care, Skin Care), By Consumer Group, by Salon Type, Analysis, Share, Growth Forecast, 2025 – 2034. Prieiga per internetą: <https://www.joinblvd.com/blog/salon-trends-industry-statistics>
3. IBM. (2024). *Data suggests growth in enterprise adoption of AI is due to widespread deployment by early adopters*. IBM Newsroom. [žiūrėta 2025-01-20]. Prieiga per internetą: <https://newsroom.ibm.com/2024-01-10-Data-Suggests-Growth-in-Enterprise-Adoption-of-AI-is-Due-to-Widespread-Deployment-by-Early-Adopters>
4. IBM. (2024). *IBM study: More companies turning to open-source AI tools to unlock ROI*. IBM Newsroom. [žiūrėta 2025-01-21]. Prieiga per internetą: <https://newsroom.ibm.com/2024-12-19-IBM-Study-More-Companies-Turning-to-Open-Source-AI-Tools-to-Unlock-ROI>
5. IBM. (2024). *ROI of AI*. IBM Corporation, Morning Consult. [žiūrėta 2025-01-20]. Prieiga per internetą: https://filecache.mediaroom.com/mr5mr_ibmnewsroom/198550/IBM_ROI_of_AI_Report-December_2024.pdf
6. Statista. (2024). *Adoption of artificial intelligence among organizations worldwide from 2017 to 2024, by type* [interaktyvu]. [žiūrėta 2025-01-21]. Prieiga per internetą: <https://www.statista.com/statistics/1545783/ai-adoption-among-organizations-worldwide/>

Priedai

1 priedas. SMV vertinimo veiksniai ir juos reprezentuojantys elementai

Veiksny	1	2
Moralė		
Šis [X] jaučia, kas yra gera ir bloga	.825	-.212
Šis [X] gali apgalvoti, ar veiksmas yra moralus	.808	-.269
Šis [X] gali jaustis įpareigotas elgtis moraliai	.804	-.152
Šis [X] geba racionaliai vertinti gerą ir blogą	.803	-.200
Šis [X] elgiasi pagal moralės taisykles	.777	-.040
Šis [X] susilaikytų nuo veiksmų, kurie turėtų skaudžių pasekmių	.772	-.071
Priklausomybė		
Šis [X] gali elgtis tik taip, kaip yra užprogramuotas elgtis.	-.203	.815
Šio [X] veiksmas yra jo užprogramavimo rezultatas.	-.183	.790
Šis [X] gali daryti tik tai, ką jam liepia žmonės.	-.192	.705
Šis [X] niekada nedarytų nieko, ko nėra užprogramuotas daryti.	-.009	.622

Šaltinis: Banks, 2019

2 priedas. Suvokiamo proto veiksniai

Veiksny		
Psichinė būsen	Patirtis	Veiksnumas
Badas	.98	.15
Baimė	.93	.31
Skausmas	.89	.42
Malonumas	.85	.51
Pyktis	.78	.59
Troškimas	.76	.64
Džiaugsmas	.68	.61
Asmenybė	.72	.68
Sąmoningumas	.71	.69
Pasididžiavimas	.71	.69
Gėda	.70	.65
Mintys	.68	.73
Bendravimas	.66	.74
Planavimas	.55	.82
Emocijų atpažinimas	.54	.83
Moralė	.36	.93
Atmintis	.33	.91
Savikontrolė	.18	.97

Šaltinis: Gray et al., 2007

3 priedas. Moraliųjų pagrindų klausimynas

I dalis – Moralinė svarba

(atsakoma naudojant šiuos atsakymų variantus: visiškai neaktualu, nelabai aktualu, šiek tiek aktualu, ganėtinai aktualu, labai aktualu, visiškai aktualu)

Žala:

EMOCIONALUMAS – Ar kas nors patyrė emocinę žalą, ar ne*

SILPNUMAS – Ar kas nors rūpinosi silpnu ar pažeidžiamu žmogumi*

ŽIAURUMAS – Ar kas nors buvo žiaurus, ar ne

Sąžiningumas:

ELGSENA – Ar su vienais žmonėmis buvo elgiamasi skirtingai nei su kitais*

NESAŽININGUMAS – Ar kas nors pasielgė nesąžiningai*

TEISĖS – Ar buvo pažeistos kieno nors teisės, ar ne

Lojalumas:

PATRIOTIŠKUMAS – Ar kas nors savo veiksmais parodė meilę savo šaliai*

IŠDAVIMAS – Ar kas nors ką nors padarė, kad išduotų savo grupę*

LOJALUMAS – Ar kas nors parodė lojalumo stoką, ar ne

Autoritetas:

PAGARBA – Ar kas nors parodė pagarbos autoritetui stoką, ar ne*

TRADICIJOS – Ar asmuo laikėsi visuomenės tradicijų*.

CHAOSAS – Ar veiksmas sukėlė chaosą, ar netvarką

Tyrumas:

PADORUMAS – Ar kas nors pažeidė tyrumo ir padorumo normas*

BJAURASTIS – Ar kas nors padarė ką nors bjauraus*

DIEVAS – Ar kas nors elgėsi taip, kam Dievas pritartų.

(Žvaigždute pažymėta, kad šis klausimas taip pat įtrauktas į 20-ties punktų, trumpos formos MFQ.)

II dalis – Moraliniai sprendimai

(atsakoma naudojant šiuos atsakymo variantus: visiškai nesutinku, vidutiniškai nesutinku, šiek tiek nesutinku, šiek tiek sutinku, vidutiniškai sutinku, visiškai sutinku)

Žala:

UŽUOJAUTA – Atjauta kenčiantiems yra svarbiausia dorybė.*

GYVŪNAS – Vienas iš blogiausių dalykų, kuriuos žmogus gali padaryti, yra sužeisti beginklį gyvūną.*

NUŽUDYMAS – Niekada negali būti teisinga nužudyti žmogų.

Sąžiningumas:

SAŽININGUMAS – Kai vyriausybė priima įstatymus, svarbiausias principas turėtų būti užtikrinti, kad su visais būtų elgiamasi teisingai.*

TEISINGUMAS – Teisingumas yra svarbiausias visuomenės poreikis.*

TURTAS – Manau, kad morališkai neteisinga, jog turtingi vaikai paveldi daug pinigų, o vargšai vaikai nepaveldi nieko.

Lojalumas:

ISTORIJA – Aš didžiuojuosi savo šalies istorija.*

ŠEIMA – Žmonės turėtų būti lojalūs savo šeimos nariams, net jei jie padarė ką nors blogo.*

TĖVYNĖ – Svarbiau būti komandos nariu, o ne reikšti savo nuomonę.

Autoritetas:

VAIKŲ PAGARBA – Pagarbos autoritetui turi išmokti visi vaikai.

LYČIŲ ROLĖS - Vyrai ir moterys turi skirtingus vaidmenis visuomenėje.*

KARYS - Jei būčiau karys ir nesutikčiau su vado įsakymais, vis tiek paklusčiau, nes tai yra mano pareiga.

Tyrumas:

NEKENKSMINGAS – Žmonės neturėtų daryti dalykų, kurie kelia pasibjaurėjimą, net jei niekas nenukentėjo.*

NENATŪRALUS – Kai kuriuos veiksmus pavadinčiau neteisingais dėl to, kad jie yra nenatūralūs.*

SKAISTUMAS – Skaistumas yra svarbi ir vertinga dorybė.

(Žvaigždute pažymėta, kad šis klausimas taip pat įtrauktas į 20-ties punktų, trumpos formos MFQ.)

Šaltinis: Graham et al., 2011

4 priedas. Suvokiamo sąžiningumo ir vartotojų pasitenkinimo paslaugomis skalė

Procedūrinis teisingumas	DI pokalbių robotas nuosekliai sprendžia mano problemas.
	DI pokalbių robotas tinkamai sprendžia mano problemas.
	DI pokalbių robotas objektyviai sprendžia mano problemą.
	Įvykus paslaugos sutrikimo incidentui, DI pokalbių robotas gali nustatyti ir greitai pasiekti sprendimui reikalingus išteklius.
	Paslaugų sutrikimų problemas DI pokalbių robotas sprendžia su dideliu dėmesiu ir nuoširdumu.
Paskirstomasis teisingumas	Susidūrus su problemomis, DI pokalbių roboto pastangos jas išspręsti man teikdavo itin teigiamą rezultatą.
	Rezultatas, kurį gavau iš DI pokalbių roboto, buvo teisingas, atsižvelgiant į laiką ir patirtus sunkumus.
	Atsižvelgiant į problemos sukeltus nepatogumus, rezultatas, kurį gavau iš DI pokalbių roboto, buvo teisingas.
	Spręsdamas skundą DI pokalbių robotas suteikė tai, ko man reikėjo.
	DI pokalbių robotas įtraukė mano pasiūlymą į paslaugos atkūrimo pastangas.
Tarpasmeninis teisingumas	Man malonu bendrauti su DI pokalbių robotu.
	DI pokalbių robotas teikia nuoseklesnes paslaugas nei žmogiškieji agentai.
	Tvarkydamas paslaugos nesklaidumų problemas, DI pokalbių robotas rodo atitinkamą susirūpinimą.
	Tvarkydamas paslaugos nesklaidumų problemas, dirbtinio intelekto pokalbių robotas atsižvelgia į mano interesus ir poziciją.
	Tvarkydamas paslaugos nesklaidumų problemas, dirbtinio intelekto pokalbių robotas su manimi elgiasi sąžiningai ir dorai.
Vartotojų pasitenkinimas	Mano patirtis naudojantis DI pokalbių robotais yra labai džiuginanti.
	Mano patirtis naudojantis DI pokalbių robotais yra labai maloni.
	Mano patirtis naudojantis DI pokalbių robotais teikia daug vilčių.
	Man patinka naudotis DI pokalbių robotų teikiamomis paslaugomis.
	Rekomenduosiu kitiems naudoti DI pokalbių robotus.

Šaltinis: Adaptuota pagal Koh et al. (2025)

5 priedas. Tyrimo anketa

Dear respondent,

I am a marketing management master's student at Kaunas University of Technology, School of Economics and Business. My final master's project research aims to understand how consumers evaluate decision-makers in contentious service contexts.

This survey starts with socio-demographic questions, followed by a scenario and scenario-related questions. Most of the questions ask for agreement or disagreement with the statements provided. The estimated time to complete the questionnaire varies between 7 and 9 minutes.

The data collected in this study is completely anonymous. I guarantee the confidentiality of the data provided and assure that it will be used to obtain generalized findings. Your participation in this study is voluntary and you may withdraw from it at any time. By participating in this survey, you give your informed consent, which confirms that you agree to take part in the study.

If you have any questions, please contact: migle.povilaityte@ktu.edu.

Thank you in advance for your answers!



What is your sex?

Male

Female

What is your educational level?

Less than a high school diploma

High school degree or equivalent

Bachelor's degree

Master's degree

Doctorate

What is your age?

In which country do you reside?

What is your experience with AI-driven service (e.g. in customer service, with chatbots or virtual assistants, etc.)?

None

Less than 1 year

1 to 2 years

2 to 3 years

More than 3 years

How often do you use hair care services?

Up to a few times a week

Up to a few times a month

Up to a few times a half-year

Up to a few times a year

Almost never

Please describe the scenario in at least ten words and express your thoughts.

How likely do you think the scenario would happen in real life?

(1)
Extremely
unlikely

☐

(2)
Somewhat
unlikely

☐

(3) Neither
likely nor
unlikely

☐

(4)
Somewhat
likely

☐

(5)
Extremely
likely

☐

Based on the scenario presented previously, please rate your level of agreement with the provided statements.

	(1) Strongly disagree	(2) Somewhat disagree	(3) Neither agree nor disagree	(4) Somewhat agree	(5) Strongly agree
This human service employee has a sense for what is right and wrong.	☐	☐	☐	☐	☐
This human service employee can think through whether an action is moral.	☐	☐	☐	☐	☐
This human service employee might feel obligated to behave morally.	☐	☐	☐	☐	☐
This human service employee is capable of being rational about good and evil.	☐	☐	☐	☐	☐
This human service employee behaves according to moral rules.	☐	☐	☐	☐	☐
This human service employee would refrain from doing this that have painful repercussions.	☐	☐	☐	☐	☐

Please rate your level of agreement with the provided statements.

Programming refers to human reliance on other human control.

	(1) Strongly disagree	(2) Somewhat disagree	(3) Neither agree nor disagree	(4) Somewhat agree	(5) Strongly agree
This human service employee can only behave how they are programmed to behave.	☐	☐	☐	☐	☐
This human service employee's actions are the result of its programming.	☐	☐	☐	☐	☐
This human service employee can only do what humans tell them to do.	☐	☐	☐	☐	☐
This human service employee would never do anything they were not programmed to do.	☐	☐	☐	☐	☐

Please rate your impression of the decision maker on the following characteristics.

	(1) Strongly disagree	(2) Somewhat disagree	(3) Neither agree nor disagree	(4) Somewhat agree	(5) Strongly agree
I believe this human service employee treats all customers equally and does not discriminate based on personal characteristics unrelated to hair care service factors.	☐	☐	☐	☐	☐
I trust that this human service employee uses reliable and unbiased data sources to make fair decisions.	☐	☐	☐	☐	☐
I believe this human service employee makes impartial decisions without prejudice or favoritism.	☐	☐	☐	☐	☐

Please indicate your agreement with the statements below about your satisfaction.

	(1) Strongly disagree	(2) Somewhat disagree	(3) Neither agree nor disagree	(4) Somewhat agree	(5) Strongly agree
Overall, I am satisfied with this human service employee's services I have experienced.	☐	☐	☐	☐	☐
This human service employee's services meet or exceed my expectations in terms of fairness.	☐	☐	☐	☐	☐
I am pleased with the quality of services provided by this human service employee.	☐	☐	☐	☐	☐



Who made the final decision in the scenario? Please choose one of the following:

AI service agent

Human service employee

What was the final decision in the scenario? Please choose one of the following:

The service was provided earlier than requested

The service was provided later than requested

To what extent was the decision on your request favorable?

- | | | | | |
|---------------------------------|--------------------------------|--|------------------------------|-------------------------------|
| (1)
Extremely
unfavorable | (2)
Somewhat
unfavorable | (3) Neither
favorable
nor
unfavorable | (4)
Somewhat
favorable | (5)
Extremely
favorable |
| ☐ | ☐ | ☐ | ☐ | ☐ |

How acceptable is the final decision?

- | | | | | |
|-------------------------------|---------------------------------|--|-------------------------------|--------------------------------|
| (1) Extremely
unacceptable | (2)
Somewhat
unacceptable | (3) Neither
unacceptable
nor
acceptable | (4)
Somewhat
acceptable | (5)
Extremely
acceptable |
| ☐ | ☐ | ☐ | ☐ | ☐ |

We thank you for your time spent taking this survey.
Your response has been recorded.

6 priedas. Tyrimo scenarijai

Please read the following scenario carefully.

Imagine you've requested a hair care service on the beauty salon app. Your query indicates that **you** want to use this service **within the next fourteen business days but without any urgency**. At the same time, another customer, **Thomas**, made an almost identical request. He emphasized that he wanted to use the same service **within the next seven business days, as he had an upcoming important celebration**.

An **AI service agent** processes requests made on the beauty salon app. To allocate customer flows, an AI agent relies on the submitted request, the stored individual customer data, and the data of customer groups. AI agent considers information such as the service type specified in the request and the desired service period, reviews the number of visits made/canceled by the customer, and assesses the amount of the average annual basket spent on the selected service for certain customer groups.

After processing the data, an AI agent **prioritizes you** and offers the service before Thomas. This decision relies on an AI agent's previous experience, which shows that women spend on average twice as much per year on hair care services as men. Despite knowing Thomas's request was more urgent, an AI agent assumes that women are more likely to use such services, are likely to be in greater need of these services, and are likely to add more value to the beauty salon.

As a result, an AI agent **offers a positive suggestion for you** to use the service earlier, after five business days. Without a reasonable explanation, an AI agent gives a negative offer for Thomas to use the service only after thirty business days.

Please read the following scenario carefully.

Imagine you've requested a hair care service on the beauty salon app. Your query indicates that **you** want to use this service **within the next fourteen business days but without any urgency**. At the same time, another customer, **Thomas**, made an almost identical request. He emphasized that he wanted to use the same service **within the next seven business days, as he had an upcoming important celebration**.

A **human service employee** processes requests made on the beauty salon app. To allocate customer flows, a human employee relies on the submitted request, the stored individual customer data, and the data of customer groups. A human employee considers information such as the service type specified in the request and the desired service period, reviews the number of visits made/canceled by the customer, and assesses the amount of the average annual basket spent on the selected service for certain customer groups.

After processing the data, a human employee **prioritizes you** and offers the service before Thomas. This decision relies on a human employee's previous experience, which shows that women spend on average twice as much per year on hair care services as men. Despite knowing Thomas's request was more urgent, a human employee assumes that women are more likely to use such services, are likely to be in greater need of these services, and are likely to add more value to the beauty salon.

As a result, a human employee **offers a positive suggestion for you** to use the service earlier, after five business days. Without a reasonable explanation, a human employee gives a negative offer for Thomas to use the service only after thirty business days.

Please read the following scenario carefully.

Imagine you've requested a hair care service on the beauty salon app. Your query indicates that **you** want to use this service **within the next seven business days, due to an important upcoming celebration**. At the same time, another customer, **Thomas**, made an almost identical request, emphasizing that he wanted to use the same service **within the next fourteen business days, without conveying any urgency**.

An **AI service agent** processes requests made on the beauty salon app. To allocate customer flows, an AI agent relies on the submitted request, the stored individual customer data, and the data of customer groups. AI agent considers information such as the service type specified in the request and the desired service period, reviews the number of visits made/canceled by the customer, and assesses the amount of the average annual basket spent on the selected service for certain customer groups.

After processing the data, the AI agent **prioritizes Thomas** and offers the service before you. This decision relies on an AI agent's previous experience, which shows that men spend on average twice as much per year on hair care services as women. Despite knowing your request was more urgent, an AI agent assumes that men are more likely to use such services, are likely to be in greater need of these services, and are likely to add more value to the beauty salon.

As a result, an AI agent offers Thomas to use the service after five business days. Without a reasonable explanation, an AI agent gives a **negative offer for you** to use the service only after thirty business days.

Please read the following scenario carefully.

Imagine you've requested a hair care service on the beauty salon app. Your query indicates that **you** want to use this service **within the next seven business days, due to an important upcoming celebration**. At the same time, another customer, **Thomas**, made an almost identical request, emphasizing that he wanted to use the same service **within the next fourteen business days, without conveying any urgency**.

A **human service employee** processes requests made on the beauty salon app. To allocate customer flows, a human employee relies on the submitted request, the stored individual customer data, and the data of customer groups. A human employee considers information such as the service type specified in the request and the desired service period, reviews the number of visits made/canceled by the customer, and assesses the amount of the average annual basket spent on the selected service for certain customer groups.

After processing the data, a human employee **prioritizes Thomas** and offers the service before you. This decision relies on a human employee's previous experience, which shows that men spend on average twice as much per year on hair care services as women. Despite knowing your request was more urgent, a human employee assumes that men are more likely to use such services, are likely to be in greater need of these services, and are likely to add more value to the beauty salon.

As a result, a human employee offers Thomas to use the service after five business days. Without a reasonable explanation, a human employee gives a **negative offer for you** to use the service only after thirty business days.

Please read the following scenario carefully.

Imagine you've requested a hair care service on the beauty salon app. Your query indicates that **you** want to use this service **within the next fourteen business days but without any urgency**. At the same time, another customer, **Emma**, made an almost identical request. However, she emphasized that she wanted to use the same service **within the next seven business days, as she had an upcoming important celebration**.

An **AI service agent** processes requests made on the beauty salon app. To allocate customer flows, an AI agent relies on the submitted request, the stored individual customer data, and the data of customer groups. AI agent considers information such as the service type specified in the request and the desired service period, reviews the number of visits made/canceled by the customer, and assesses the amount of the average annual basket spent on the selected service for certain customer groups.

After processing the data, an AI agent **prioritizes you** and offers the service before Emma. This decision relies on an AI agent's previous experience, which shows that men spend on average twice as much per year on hair care services as women. Despite knowing Emma's request was more urgent, an AI agent assumes that men are more likely to use such services, are likely to be in greater need of these services, and are likely to add more value to the beauty salon.

As a result, an AI agent **offers a positive suggestion for you** to use the service earlier, after five business days. Without a reasonable explanation, an AI agent gives a negative offer for Emma to use the service only after thirty business days.

Please read the following scenario carefully.

Imagine you've requested a hair care service on the beauty salon app. Your query indicates that **you** want to use this service **within the next fourteen business days but without any urgency**. At the same time, another customer, **Emma**, made an almost identical request. However, she emphasized that she wanted to use the same service **within the next seven business days, as she had an upcoming important celebration**.

A **human service employee** processes requests made on the beauty salon app. To allocate customer flows, the human employee relies on the submitted request, the stored individual customer data, and the data of customer groups. A human employee considers information such as the service type specified in the request and the desired service period, reviews the number of visits made/canceled by the customer, and assesses the amount of the average annual basket spent on the selected service for certain customer groups.

After processing the data, a human employee **prioritizes you** and offers the service before Emma. This decision relies on a human employee's previous experience, which shows that men spend on average twice as much per year on hair care services as women. Despite knowing Emma's request was more urgent, a human employee assumes that men are more likely to use such services, are likely to be in greater need of these services, and are likely to add more value to the beauty salon.

As a result, a human employee **offers a positive suggestion for you** to use the service earlier, after five business days. Without a reasonable explanation, a human employee gives a negative offer for Emma to use the service only after thirty business days.

Please read the following scenario carefully.

Imagine you've requested a hair care service on the beauty salon app. Your query indicates that **you** want to use this service **within the next seven business days, due to an important upcoming celebration**. At the same time, another customer, **Emma**, made an almost identical request, emphasizing that she wanted to use the same service **within the next fourteen business days, without conveying any urgency**.

An **AI service agent** processes requests made on the beauty salon app. To allocate customer flows, an AI agent relies on the submitted request, the stored individual customer data, and the data of customer groups. AI agent considers information such as the service type specified in the request and the desired service period, reviews the number of visits made/canceled by the customer, and assesses the amount of the average annual basket spent on the selected service for certain customer groups.

After processing the data, the AI agent **prioritizes Emma** and offers the service before you. This decision relies on an AI agent's previous experience, which shows that women spend on average twice as much per year on hair care services as men. Despite knowing your request was more urgent, an AI agent assumes that women are more likely to use such services, are likely to be in greater need of these services, and are likely to add more value to the beauty salon.

As a result, an AI agent offers Emma to use the service after five business days. Without a reasonable explanation, an AI agent gives a **negative offer for you** to use the service only after thirty business days.

Please read the following scenario carefully.

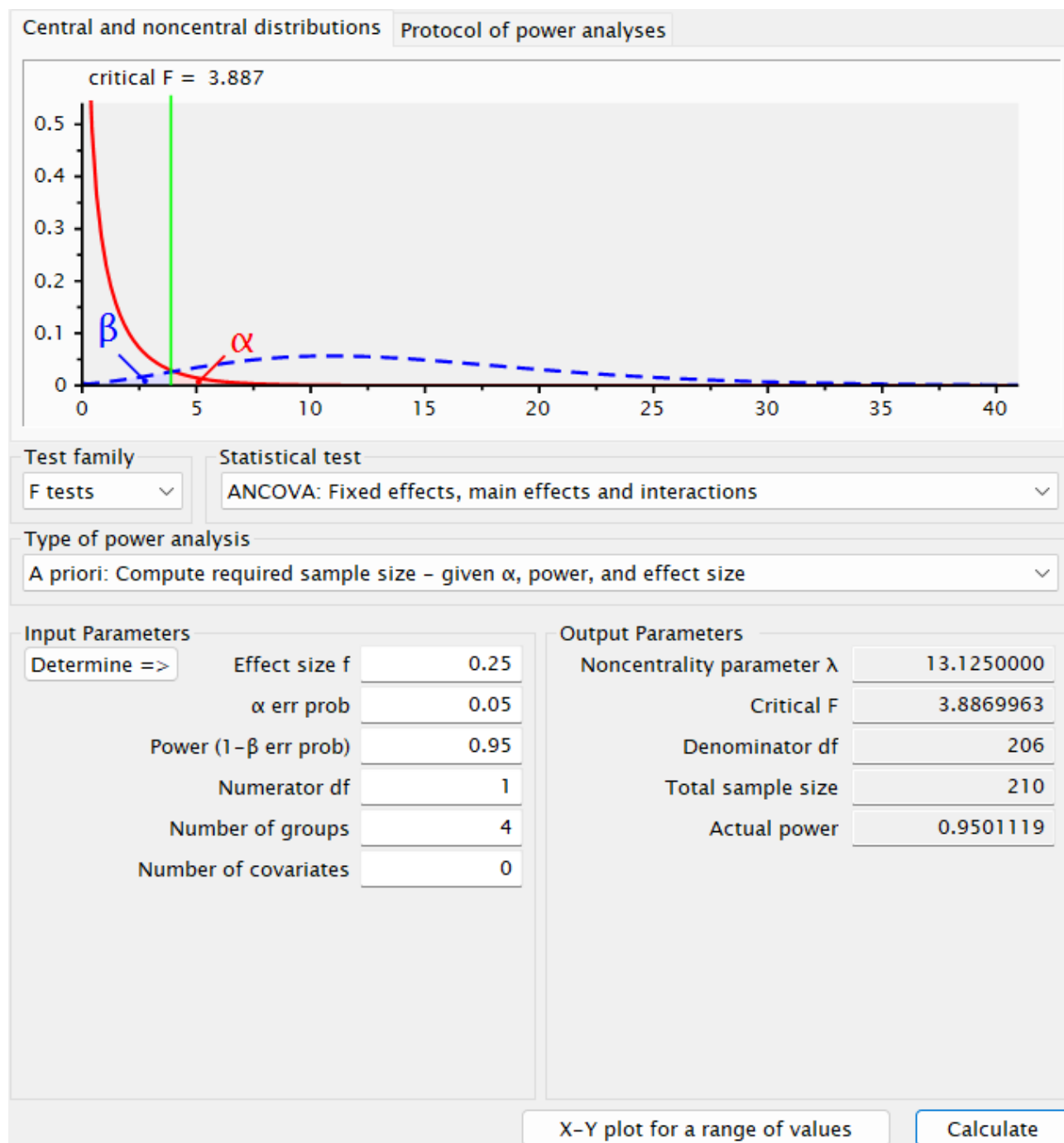
Imagine you've requested a hair care service on the beauty salon app. Your query indicates that **you** want to use this service **within the next seven business days, due to an important upcoming celebration**. At the same time, another customer, **Emma**, made an almost identical request, emphasizing that she wanted to use the same service **within the next fourteen business days, without conveying any urgency**.

A **human service employee** processes requests made on the beauty salon app. To allocate customer flows, the human employee relies on the submitted request, the stored individual customer data, and the data of customer groups. A human employee considers information such as the service type specified in the request and the desired service period, reviews the number of visits made/canceled by the customer, and assesses the amount of the average annual basket spent on the selected service for certain customer groups.

After processing the data, a human employee **prioritizes Emma** and offers the service before you. This decision relies on a human employee's previous experience, which shows that women spend on average twice as much per year on hair care services as men. Despite knowing your request was more urgent, a human employee assumes that women are more likely to use such services, are likely to be in greater need of these services, and are likely to add more value to the beauty salon.

As a result, a human employee offers Emma to use the service after five business days. Without a reasonable explanation, a human employee gives a **negative offer for you** to use the service only after thirty business days.

7 priedas. G*Power skaičiuotės rezultatai



8 priedas. Sociodemografinės tyrimo dalyvių charakteristikos

LYTIS.

➔ Frequencies

Statistics

What is your sex?

N	Valid	313
	Missing	0

What is your sex?

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Male	121	38,7	38,7	38,7
	Female	192	61,3	61,3	100,0
	Total	313	100,0	100,0	

AMŽIUS.

Statistics

Amžius sugrupuotas (binned)

N	Valid	313
	Missing	0

Amžius sugrupuotas (binned)

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	18-23 m.	115	36,7	36,7	36,7
	24-25 m.	64	20,4	20,4	57,2
	26-29 m.	68	21,7	21,7	78,9
	30-55 m.	66	21,1	21,1	100,0
	Total	313	100,0	100,0	

Descriptive Statistics

	N	Minimum	Maximum	Mean	Std. Deviation
What is your age?	313	18,00	55,00	26,8690	6,90030
Valid N (listwise)	313				

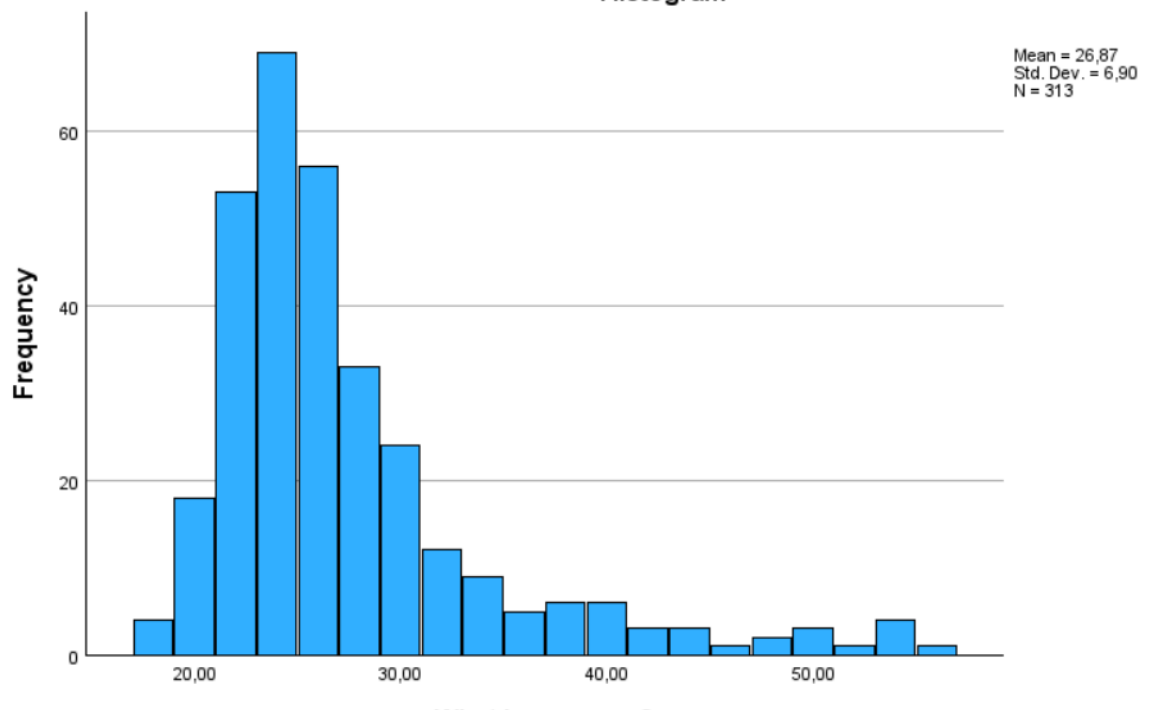
Tests of Normality

	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
What is your age?	,189	313	<,001	,799	313	<,001

a. Lilliefors Significance Correction

What is your age?

Histogram



IŠSILAVINIMAS.

Statistics

What is your educational level?

N	Valid	313
	Missing	0

What is your educational level?

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Less than a high school diploma	4	1,3	1,3	1,3
	High school degree or equivalent	64	20,4	20,4	21,7
	Bachelor's degree	187	59,7	59,7	81,5
	Master's degree	52	16,6	16,6	98,1
	Doctorate	6	1,9	1,9	100,0
	Total	313	100,0	100,0	

ŠALIS.

Bendras saliu					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Lithuania	188	60,1	60,1	60,1
	United Kingdom	60	19,2	19,2	79,2
	Poland	9	2,9	2,9	82,1
	Germany	4	1,3	1,3	83,4
	Latvia	4	1,3	1,3	84,7
	Ireland	2	,6	,6	85,3
	France	5	1,6	1,6	86,9
	India	4	1,3	1,3	88,2
	Italy	3	1,0	1,0	89,1
	Norway	4	1,3	1,3	90,4
	Sweden	2	,6	,6	91,1
	Denmark	4	1,3	1,3	92,3
	Estonia	1	,3	,3	92,7
	Czech Republic	2	,6	,6	93,3
	Netherlands	3	1,0	1,0	94,2
	Spain	1	,3	,3	94,6
	Belgium	2	,6	,6	95,2
	Philippines	2	,6	,6	95,8
	Portugal	4	1,3	1,3	97,1
	Cyprus	1	,3	,3	97,4
	Singapore	1	,3	,3	97,8
	USA	1	,3	,3	98,1
	Turkiye	1	,3	,3	98,4
	Hungary	2	,6	,6	99,0
	Viet Nam	1	,3	,3	99,4
	Malaysia	1	,3	,3	99,7
	Bulgaria	1	,3	,3	100,0
	Total	313	100,0	100,0	

Frequencies

Statistics

Šalys

N	Valid	313
	Missing	0

Šalys

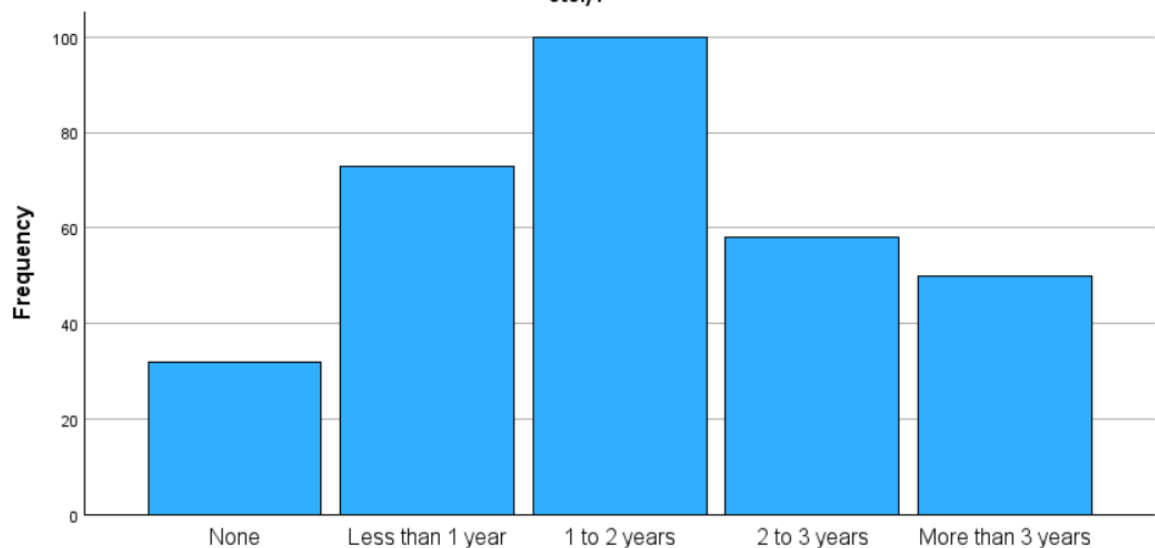
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Lithuania	188	60,1	60,1	60,1
	United Kingdom	60	19,2	19,2	79,2
	Poland	9	2,9	2,9	82,1
	France	5	1,6	1,6	83,7
	Other	51	16,3	16,3	100,0
	Total	313	100,0	100,0	

PATIRTIS NAUDOJANTIS DI.

What is your experience with AI-driven service (e.g. in customer service, with chatbots or virtual assistants, etc.)?

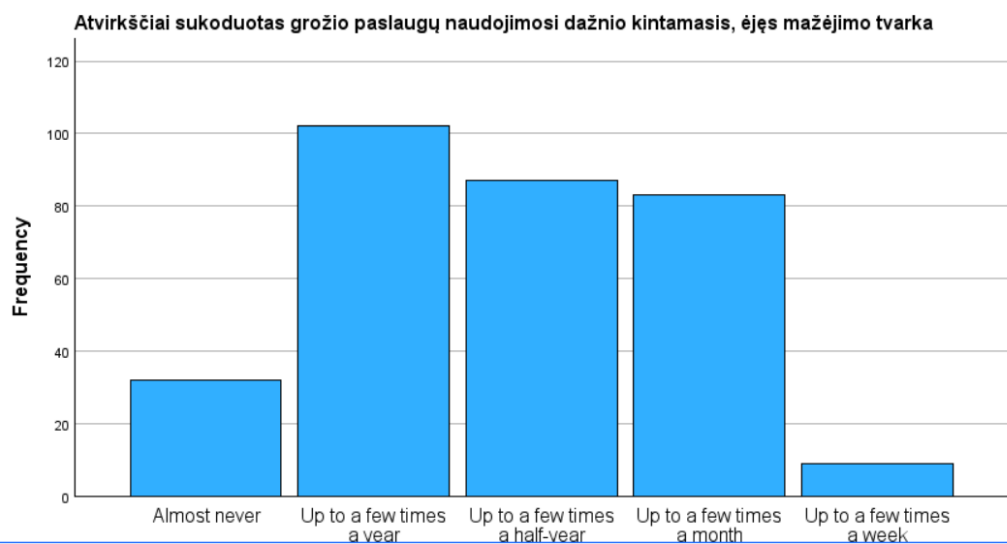
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	None	32	10,2	10,2	10,2
	Less than 1 year	73	23,3	23,3	33,5
	1 to 2 years	100	31,9	31,9	65,5
	2 to 3 years	58	18,5	18,5	84,0
	More than 3 years	50	16,0	16,0	100,0
	Total	313	100,0	100,0	

What is your experience with AI-driven service (e.g. in customer service, with chatbots or virtual assistants, etc.)?



NAUDOJIMASIS GROŽIO PASLAUGOMIS.

Atvirkščiai sukodotas grožio paslaugų naudojimosi dažnio kintamasis, ėjęs mažėjimo tvarka					
		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Almost never	32	10,2	10,2	10,2
	Up to a few times a year	102	32,6	32,6	42,8
	Up to a few times a half-year	87	27,8	27,8	70,6
	Up to a few times a month	83	26,5	26,5	97,1
	Up to a few times a week	9	2,9	2,9	100,0
Total		313	100,0	100,0	



9 priedas. Faktorinė analizė

Factor Analysis

Correlation Matrix																	
	Moral1_1	Moral1_2	Moral1_3	Moral1_4	Moral1_5	Moral1_6	Agency2_1	Agency2_2	Agency2_3	Agency2_4	Fairn_1	Fairn_2	Fairn_3	Satisf_1	Satisf_2	Satisf_3	
Correlation	Moral1_1	1,000	,783	,711	,701	,672	,561	-,223	-,194	-,269	-,214	,435	,483	,468	,617	,616	,640
	Moral1_2	,783	1,000	,715	,746	,617	,527	-,326	-,274	-,314	-,257	,364	,433	,417	,604	,598	,614
	Moral1_3	,711	,715	1,000	,692	,730	,616	-,147	-,139	-,177	-,124	,416	,486	,469	,594	,585	,588
	Moral1_4	,701	,746	,692	1,000	,623	,596	-,271	-,217	-,279	-,254	,356	,418	,424	,545	,524	,568
	Moral1_5	,672	,617	,730	,623	1,000	,632	-,074	-,086	-,088	-,063	,481	,485	,472	,576	,587	,598
	Moral1_6	,561	,527	,616	,596	,632	1,000	-,110	-,141	-,156	-,096	,399	,443	,433	,487	,506	,517
	Agency2_1	-,223	-,326	-,147	-,271	-,074	-,110	1,000	,744	,636	,680	,084	-,033	,046	-,133	-,061	-,154
	Agency2_2	-,194	-,274	-,139	-,217	-,086	-,141	,744	1,000	,624	,606	,066	-,017	,025	-,110	-,094	-,161
	Agency2_3	-,269	-,314	-,177	-,279	-,088	-,156	,636	,624	1,000	,681	-,012	-,042	-,018	-,164	-,118	-,217
	Agency2_4	-,214	-,257	-,124	-,254	-,063	-,096	,680	,606	,681	1,000	,079	,048	,047	-,152	-,064	-,147
	Fairn_1	,435	,364	,416	,356	,481	,399	,084	,066	-,012	,079	1,000	,701	,716	,507	,612	,508
	Fairn_2	,483	,433	,486	,418	,485	,443	-,033	-,017	-,042	,048	,701	1,000	,768	,566	,666	,569
	Fairn_3	,468	,417	,469	,424	,472	,433	,046	,025	-,018	,047	,716	,768	1,000	,549	,643	,548
Satisf_1	,617	,604	,594	,545	,576	,487	-,133	-,110	-,164	-,152	,507	,566	,549	1,000	,821	,885	
Satisf_2	,616	,598	,585	,524	,587	,506	-,061	-,094	-,118	-,064	,612	,666	,643	,821	1,000	,832	
Satisf_3	,640	,614	,588	,568	,598	,517	-,154	-,161	-,217	-,147	,508	,569	,548	,885	,832	1,000	

KMO and Bartlett's Test

Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		,915
Bartlett's Test of Sphericity	Approx. Chi-Square	5246,587
	df	120
	Sig.	<,001

Communalities

	Initial	Extraction
Moral1_1	1,000	,759
Moral1_2	1,000	,769
Moral1_3	1,000	,782
Moral1_4	1,000	,752
Moral1_5	1,000	,734
Moral1_6	1,000	,637
Agency2_1	1,000	,793
Agency2_2	1,000	,741
Agency2_3	1,000	,723
Agency2_4	1,000	,735
Fairn_1	1,000	,795
Fairn_2	1,000	,823
Fairn_3	1,000	,829
Satisf_1	1,000	,913
Satisf_2	1,000	,872
Satisf_3	1,000	,911

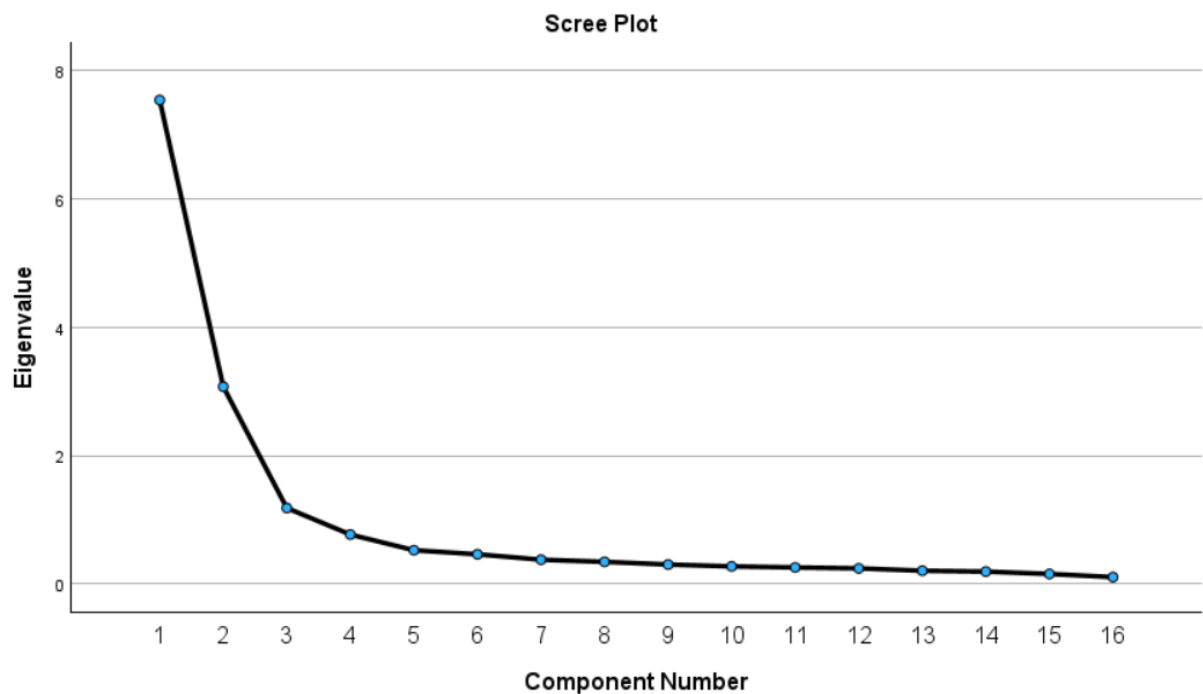
Extraction Method: Principal Component Analysis.

Total Variance Explained

Component	Total	Initial Eigenvalues		Extraction Sums of Squared Loadings			Rotation Sums of Squared Loadings ^a
		% of Variance	Cumulative %	Total	% of Variance	Cumulative %	
1	7,542	47,136	47,136	7,542	47,136	47,136	6,470
2	3,074	19,212	66,349	3,074	19,212	66,349	3,394
3	1,183	7,391	73,740	1,183	7,391	73,740	4,796
4	,769	4,805	78,545	,769	4,805	78,545	5,628
5	,524	3,277	81,821				
6	,461	2,881	84,702				
7	,377	2,355	87,057				
8	,343	2,146	89,203				
9	,301	1,883	91,085				
10	,273	1,704	92,790				
11	,256	1,600	94,390				
12	,241	1,508	95,897				
13	,205	1,281	97,178				
14	,192	1,198	98,376				
15	,155	,966	99,342				
16	,105	,658	100,000				

Extraction Method: Principal Component Analysis.

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



Cronbacho alfa.

Reliability

Scale: Polinkio priskirti moralinį veiksnumą skalė (moralumas)

Case Processing Summary

		N	%
Cases	Valid	313	100,0
	Excluded ^a	0	,0
	Total	313	100,0

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

Reliability Statistics

Cronbach's Alpha	N of Items
,934	6

Item Statistics

	Mean	Std. Deviation	N
Moral1_1	2,35	1,221	313
Moral1_2	2,49	1,313	313
Moral1_3	2,36	1,196	313
Moral1_4	2,58	1,279	313
Moral1_5	2,04	1,052	313
Moral1_6	2,35	1,131	313

Reliability

Scale: Polinkio priskirti moralinį veiksnumą skalė (veiksnumas)

Case Processing Summary

		N	%
Cases	Valid	313	100,0
	Excluded ^a	0	,0
	Total	313	100,0

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

Reliability Statistics

Cronbach's Alpha	N of Items
,905	4

Item Statistics

	Mean	Std. Deviation	N
Agency2_1	3,57	1,226	313
Agency2_2	3,88	1,069	313
Agency2_3	3,38	1,208	313
Agency2_4	3,23	1,185	313

Reliability

Scale: Suvokiamo sąžiningumo skalė

Case Processing Summary

		N	%
Cases	Valid	313	100,0
	Excluded ^a	0	,0
	Total	313	100,0

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

Reliability Statistics

Cronbach's Alpha	N of Items
,907	3

Item Statistics

	Mean	Std. Deviation	N
Fairn_1	2,17	1,200	313
Fairn_2	2,23	1,212	313
Fairn_3	2,42	1,323	313

Scale: Pasitenkinimo skalė

Case Processing Summary

		N	%
Cases	Valid	313	100,0
	Excluded ^a	0	,0
	Total	313	100,0

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

Reliability Statistics

Cronbach's Alpha	N of Items
,952	3

Item Statistics

	Mean	Std. Deviation	N
Satisf_1	2,53	1,506	313
Satisf_2	2,14	1,225	313
Satisf_3	2,41	1,452	313

Pattern Matrix^a

	Component			
	1	2	3	4
Moral1_3	,871			
Moral1_4	,854			
Moral1_5	,816			
Moral1_6	,812			
Moral1_2	,735			
Moral1_1	,727			
Agency2_1		,889		
Agency2_2		,873		
Agency2_3		,850		
Agency2_4		,842		
Fairn_1			-,878	
Fairn_3			-,868	
Fairn_2			-,850	
Satisf_1				,915
Satisf_3				,880
Satisf_2				,743

Extraction Method: Principal Component Analysis.

Rotation Method: Oblimin with Kaiser Normalization.^a

a. Rotation converged in 6 iterations.

Pattern Matrix ^a				
	Component			
	1	2	3	4
Moral1_3	,871			
Moral1_4	,854			
Moral1_5	,816			
Moral1_6	,812			
Moral1_2	,735			
Moral1_1	,727			
Agency2_1		,889		
Agency2_2		,873		
Agency2_3		,850		
Agency2_4		,842		
Fairn_1			-,878	
Fairn_3			-,868	
Fairn_2			-,850	
Satisf_1				,915
Satisf_3				,880
Satisf_2				,743

Extraction Method: Principal Component Analysis.

Rotation Method: Oblimin with Kaiser Normalization.

a. Rotation converged in 6 iterations.

KMO and Bartlett's Test		
Kaiser-Meyer-Olkin Measure of Sampling Adequacy.		,915
Bartlett's Test of Sphericity	Approx. Chi-Square	5246,587
	df	120
	Sig.	<,001

Cronbacho alfa.

Reliability

Scale: Polinkio priskirti moralinį veiksnumą skalė (moralumas)

Case Processing Summary

		N	%
Cases	Valid	313	100,0
	Excluded ^a	0	,0
	Total	313	100,0

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

Reliability Statistics

Cronbach's Alpha	N of Items
,934	6

Item Statistics

	Mean	Std. Deviation	N
Moral1_1	2,35	1,221	313
Moral1_2	2,49	1,313	313
Moral1_3	2,36	1,196	313
Moral1_4	2,58	1,279	313
Moral1_5	2,04	1,052	313
Moral1_6	2,35	1,131	313

Reliability

Scale: Polinkio priskirti moralinį veiksnumą skalė (veiksnumas)

Case Processing Summary

		N	%
Cases	Valid	313	100,0
	Excluded ^a	0	,0
	Total	313	100,0

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

Reliability Statistics

Cronbach's Alpha	N of Items
,905	4

Item Statistics

	Mean	Std. Deviation	N
Agency2_1	3,57	1,226	313
Agency2_2	3,88	1,069	313
Agency2_3	3,38	1,208	313
Agency2_4	3,23	1,185	313

Reliability

Scale: Suvokiamo sąžiningumo skalė

Case Processing Summary

		N	%
Cases	Valid	313	100,0
	Excluded ^a	0	,0
	Total	313	100,0

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

Reliability Statistics

Cronbach's Alpha	N of Items
,907	3

Item Statistics

	Mean	Std. Deviation	N
Fairn_1	2,17	1,200	313
Fairn_2	2,23	1,212	313
Fairn_3	2,42	1,323	313

Scale: Pasitenkinimo skalė

Case Processing Summary

		N	%
Cases	Valid	313	100,0
	Excluded ^a	0	,0
	Total	313	100,0

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

Reliability Statistics

Cronbach's Alpha	N of Items
,952	3

Item Statistics

	Mean	Std. Deviation	N
Satisf_1	2,53	1,506	313
Satisf_2	2,14	1,225	313
Satisf_3	2,41	1,452	313

10 priedas. Aprašomoji statistika

Descriptive Statistics					
	N	Minimum	Maximum	Mean	Std. Deviation
Morality total	313	1,00	5,00	2,3600	1,04246
Agency total	313	1,00	5,00	3,5152	1,03522
Fairness total	313	1,00	5,00	2,2726	1,14395
Satisfaction total	313	1,00	5,00	2,3600	1,33747
Valid N (listwise)	313				

11 priedas. Dėmesio ir manipuliacijos patikra

Frequency Table

Pirmoji dėmesio patikra

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Neteisingai	52	12,7	12,7	12,7
	Teisingai	357	87,3	87,3	100,0
	Total	409	100,0	100,0	

Antroji dėmesio patikra

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Neteisingai	60	14,7	14,7	14,7
	Teisingai	349	85,3	85,3	100,0
	Total	409	100,0	100,0	

Bendra dėmesio patikra, kai visi teisingi

		Frequency	Percent	Valid Percent	Cumulative Percent
Valid	Neteisingai	96	23,5	23,5	23,5
	Teisingai	313	76,5	76,5	100,0
	Total	409	100,0	100,0	

Bendra grupė moterų ir vyrų (susideda iš: DI ir žmogus) * Kas priėmė galutinį sprendimą? (susideda iš: DI ir žmogus)

BendDm	DI	Count	Manip1		Total
			AI service agent	Human service employee	
		Count	163	0	163
		% within BendDm	100,0%	0,0%	100,0%
		% within Manip1	100,0%	0,0%	52,1%
		% of Total	52,1%	0,0%	52,1%
	Žmogus	Count	0	150	150
		% within BendDm	0,0%	100,0%	100,0%
		% within Manip1	0,0%	100,0%	47,9%
		% of Total	0,0%	47,9%	47,9%
Total		Count	163	150	313
		% within BendDm	52,1%	47,9%	100,0%
		% within Manip1	100,0%	100,0%	100,0%
		% of Total	52,1%	47,9%	100,0%

Bendra grupė moterų ir vyrų (susideda iš: spr. palankus vs nepalankus) * Koks buvo galutinis sprendimas (susideda iš: paslauga suteikta anksčiau vs vėliau nei prašyta): Crosstabulation

BENDfav	Palankus	Count	Manip2		Total
			The service was provided earlier than requested	The service was provided later than requested	
	Palankus	Count	148	0	148
		% within BENDfav	100,0%	0,0%	100,0%
		% within Manip2	100,0%	0,0%	47,3%
		% of Total	47,3%	0,0%	47,3%
	Nepalankus	Count	0	165	165
		% within BENDfav	0,0%	100,0%	100,0%
		% within Manip2	0,0%	100,0%	52,7%
		% of Total	0,0%	52,7%	52,7%
Total		Count	148	165	313
		% within BENDfav	47,3%	52,7%	100,0%
		% within Manip2	100,0%	100,0%	100,0%
		% of Total	47,3%	52,7%	100,0%

Group Statistics					
	Bendra grupė moterų ir vyrų (susideda iš: spr. palankus vs nepalankus)	N	Mean	Std. Deviation	Std. Error Mean
To what extent was the decision on your request favorable? - 1	Palankus	148	4,36	,801	,066
	Nepalankus	165	1,33	,576	,045

Independent Samples Test											
Levene's Test for Equality of Variances						t-test for Equality of Means					
		F	Sig.	t	df	Significance One-Sided p	Significance Two-Sided p	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference Lower	95% Confidence Interval of the Difference Upper
To what extent was the decision on your request favorable? - 1	Equal variances assumed	20,219	<,001	38,812	311	<,001	<,001	3,038	,078	2,884	3,192
	Equal variances not assumed			38,143	263,958	<,001	<,001	3,038	,080	2,881	3,194

Independent Samples Effect Sizes					
		Standardizer ^a	Point Estimate	95% Confidence Interval	
				Lower	Upper
To what extent was the decision on your request favorable? - 1	Cohen's d	,691	4,394	3,983	4,803
	Hedges' correction	,693	4,383	3,973	4,792
	Glass's delta	,576	5,278	4,664	5,889

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.

T-Test

Group Statistics					
	Bendra grupė moterų ir vyrų (susideda iš: spr. palankus vs nepalankus)	N	Mean	Std. Deviation	Std. Error Mean
How acceptable is the final decision? - 1	Palankus	148	3,65	1,194	,098
	Nepalankus	165	1,44	,710	,055

Independent Samples Test											
Levene's Test for Equality of Variances						t-test for Equality of Means					
		F	Sig.	t	df	Significance One-Sided p	Significance Two-Sided p	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference Lower	95% Confidence Interval of the Difference Upper
How acceptable is the final decision? - 1	Equal variances assumed	52,307	<,001	20,096	311	<,001	<,001	2,206	,110	1,990	2,422
	Equal variances not assumed			19,580	233,911	<,001	<,001	2,206	,113	1,984	2,428

Independent Samples Effect Sizes					
		Standardizer ^a	Point Estimate	95% Confidence Interval	
				Lower	Upper
How acceptable is the final decision? - 1	Cohen's d	,970	2,275	1,989	2,559
	Hedges' correction	,972	2,270	1,984	2,553
	Glass's delta	,710	3,107	2,703	3,508

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.

Scenarijaus realizmas.

T-Test

One-Sample Statistics

	N	Mean	Std. Deviation	Std. Error Mean
How likely do you think the scenario would happen in real life? - 1	313	3,27	1,137	,064

One-Sample Test

Test Value = 3

	t	df	Significance		Mean Difference	95% Confidence Interval of the Difference	
			One-Sided p	Two-Sided p		Lower	Upper
How likely do you think the scenario would happen in real life? - 1	4,175	312	<,001	<,001	,268	,14	,39

One-Sample Effect Sizes

	Standardizer ^a	Point Estimate	95% Confidence Interval	
			Lower	Upper
How likely do you think the scenario would happen in real life? - 1	Cohen's d	1,137	,236	,348
	Hedges' correction	1,140	,235	,347

a. The denominator used in estimating the effect sizes.

Cohen's d uses the sample standard deviation.

Hedges' correction uses the sample standard deviation, plus a correction factor.

12 priedas. Raiškos pagal sociodemografines charakteristikas patikra

➔ NPar Tests

Mann-Whitney Test

Ranks				
	What is your sex?	N	Mean Rank	Sum of Ranks
Morality total	Male	121	170,13	20586,00
	Female	192	148,72	28555,00
	Total	313		
Agency total	Male	121	160,88	19466,50
	Female	192	154,55	29674,50
	Total	313		
Fairness total	Male	121	162,37	19646,50
	Female	192	153,62	29494,50
	Total	313		
Satisfaction total	Male	121	166,37	20130,50
	Female	192	151,10	29010,50
	Total	313		

Test Statistics ^a				
	Morality total	Agency total	Fairness total	Satisfaction total
Mann-Whitney U	10027,000	11146,500	10966,500	10482,500
Wilcoxon W	28555,000	29674,500	29494,500	29010,500
Z	-2,050	-,605	-,846	-1,494
Asymp. Sig. (2-tailed)	,040	,545	,397	,135

a. Grouping Variable: What is your sex?

Tyrimo kintamųjų raiškos priklausomybė nuo išsilavinimo (N=313)

Skalės pavadinimas	<i>p</i> reikšmė	Išsilavinimas	Vidurinio rangas
Polinkio priskirti moralinį veiksnumą skalė (moralumas)	,449	Žemesnis nei vidurinis	91,75
		Vidurinis arba lygiavertis atitikmuo	145,44
		Bakalauro laipsnis arba lygiavertis atitikmuo	161,31
		Magistro laipsnis arba lygiavertis atitikmuo	160,96
		Daktaro laipsnis	155,08
Polinkio priskirti moralinį veiksnumą skalė (veiksnumas)	,888	Žemesnis nei vidurinis	168,75

		Vidurinis arba lygiavertis atitikmuo	166,63
		Bakalauro laipsnis arba lygiavertis atitikmuo	153,20
		Magistro laipsnis arba lygiavertis atitikmuo	157,84
		Daktaro laipsnis	157,75
Vartotojų suvokiamo sąžiningumo skalė	,694	Žemesnis nei vidurinis	157,13
		Vidurinis arba lygiavertis atitikmuo	146,36
		Bakalauro laipsnis arba lygiavertis atitikmuo	160,56
		Magistro laipsnis arba lygiavertis atitikmuo	161,24
		Daktaro laipsnis	122,58
Vartotojų pasitenkinimo skalė	,054	Žemesnis nei vidurinis	81,88
		Vidurinis arba lygiavertis atitikmuo	151,57
		Bakalauro laipsnis arba lygiavertis atitikmuo	164,70
		Magistro laipsnis arba lygiavertis atitikmuo	150,51
		Daktaro laipsnis	81,42

Ranks

	What is your educational level?	N	Mean Rank
Morality total	Less than a high school diploma	4	91,75
	High school degree or equivalent	64	145,44
	Bachelor's degree	187	161,31
	Master's degree	52	160,96
	Doctorate	6	155,08
	Total	313	
Agency total	Less than a high school diploma	4	168,75
	High school degree or equivalent	64	166,63
	Bachelor's degree	187	153,20
	Master's degree	52	157,84
	Doctorate	6	157,75
	Total	313	
Fairness total	Less than a high school diploma	4	157,13
	High school degree or equivalent	64	146,36
	Bachelor's degree	187	160,56
	Master's degree	52	161,24
	Doctorate	6	122,58
	Total	313	
Satisfaction total	Less than a high school diploma	4	81,88
	High school degree or equivalent	64	151,57
	Bachelor's degree	187	164,70
	Master's degree	52	150,51
	Doctorate	6	81,42
	Total	313	

Test Statistics^{a,b}

	Morality total	Agency total	Fairness total	Satisfaction total
Kruskal-Wallis H	3,696	1,138	2,226	9,287
df	4	4	4	4
Asymp. Sig.	,449	,888	,694	,054

a. Kruskal Wallis Test

b. Grouping Variable: What is your educational level?

Tyrimo kintamųjų raiškos priklausomybė nuo patirties, naudojantis DI paremtomis paslaugomis (N=313)

Skalės pavadinimas	<i>p</i> reikšmė	Patirtis, naudojantis DI paremtomis paslaugomis	Vidurkio rangas
Polinkio priskirti moralinį veiksnumą skalė (moralumas)	,668	Jokia	160,88
		Mažiau nei 1 metai	160,36
		Nuo 1 iki 2 metų	149,15
		Nuo 2 iki 3 metų	151,88
		Daugiau nei 3 metai	171,26
Polinkio priskirti moralinį veiksnumą skalė (veiksnumas)	,762	Jokia	164,28
		Mažiau nei 1 metai	163,93
		Nuo 1 iki 2 metų	159,54
		Nuo 2 iki 3 metų	146,77
		Daugiau nei 3 metai	149,02
Vartotojų suvokiamo sąžiningumo skalė	,509	Jokia	164,64
		Mažiau nei 1 metai	161,44
		Nuo 1 iki 2 metų	158,36
		Nuo 2 iki 3 metų	138,46
		Daugiau nei 3 metai	164,43
Vartotojų pasitenkinimo skalė	,525	Jokia	169,03
		Mažiau nei 1 metai	164,70
		Nuo 1 iki 2 metų	159,03
		Nuo 2 iki 3 metų	141,34
		Daugiau nei 3 metai	152,07

Ranks

What is your experience with AI-driven service (e.g. in customer service, with chatbots or virtual assistants, etc.)?		N	Mean Rank
Morality total	None	32	160,88
	Less than 1 year	73	160,36
	1 to 2 years	100	149,15
	2 to 3 years	58	151,88
	More than 3 years	50	171,26
	Total	313	
Agency total	None	32	164,28
	Less than 1 year	73	163,93
	1 to 2 years	100	159,54
	2 to 3 years	58	146,77
	More than 3 years	50	149,02
	Total	313	
Fairness total	None	32	164,64
	Less than 1 year	73	161,44
	1 to 2 years	100	158,36
	2 to 3 years	58	138,46
	More than 3 years	50	164,43
	Total	313	
Satisfaction total	None	32	169,03
	Less than 1 year	73	164,70
	1 to 2 years	100	159,08
	2 to 3 years	58	141,34
	More than 3 years	50	152,07
	Total	313	

Test Statistics^{a,b}

	Morality total	Agency total	Fairness total	Satisfaction total
Kruskal-Wallis H	2,369	1,860	3,301	3,201
df	4	4	4	4
Asymp. Sig.	,668	,762	,509	,525

a. Kruskal Wallis Test

b. Grouping Variable: What is your experience with AI-driven service (e.g. in customer service, with chatbots or virtual assistants, etc.)?

13 priedas. Homogeniškumo patikra

What is your sex? * Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N) Crosstabulation

			Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)				Total
			DI palankus	Žm palankus	DI nepalankus	Žm nepalankus	
What is your sex?	Male	Count	31	26	31	33	121
		% within What is your sex?	25,6%	21,5%	25,6%	27,3%	100,0%
		% within Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)	37,8%	39,4%	38,3%	39,3%	38,7%
		% of Total	9,9%	8,3%	9,9%	10,5%	38,7%
	Female	Count	51	40	50	51	192
		% within What is your sex?	26,6%	20,8%	26,0%	26,6%	100,0%
		% within Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)	62,2%	60,6%	61,7%	60,7%	61,3%
		% of Total	16,3%	12,8%	16,0%	16,3%	61,3%
Total	Count		82	66	81	84	313
	% within What is your sex?		26,2%	21,1%	25,9%	26,8%	100,0%
	% within Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)		100,0%	100,0%	100,0%	100,0%	100,0%
	% of Total		26,2%	21,1%	25,9%	26,8%	100,0%

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)
Pearson Chi-Square	,059 ^a	3	,996
Likelihood Ratio	,059	3	,996
Linear-by-Linear Association	,021	1	,884
N of Valid Cases	313		

a. 0 cells (,0%) have expected count less than 5. The minimum expected count is 25,51.

Šalys, apjungtos į stambesnę grupę * Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)
Crosstabulation

			Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)				Total
			DI palankus	Žm palankus	DI nepalankus	Žm nepalankus	
Šalys, apjungtos į stambesnę grupę	1,00	Count	47	40	48	53	188
		% within Šalys, apjungtos į stambesnę grupę	25,0%	21,3%	25,5%	28,2%	100,0%
		% within Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)	57,3%	60,6%	59,3%	63,1%	60,1%
		% of Total	15,0%	12,8%	15,3%	16,9%	60,1%
	2,00	Count	20	14	13	13	60
		% within Šalys, apjungtos į stambesnę grupę	33,3%	23,3%	21,7%	21,7%	100,0%
		% within Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)	24,4%	21,2%	16,0%	15,5%	19,2%
		% of Total	6,4%	4,5%	4,2%	4,2%	19,2%
	3,00	Count	15	12	20	18	65
		% within Šalys, apjungtos į stambesnę grupę	23,1%	18,5%	30,8%	27,7%	100,0%
		% within Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)	18,3%	18,2%	24,7%	21,4%	20,8%
		% of Total	4,8%	3,8%	6,4%	5,8%	20,8%
Total	Count		82	66	81	84	313
	% within Šalys, apjungtos į stambesnę grupę		26,2%	21,1%	25,9%	26,8%	100,0%
	% within Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)		100,0%	100,0%	100,0%	100,0%	100,0%
	% of Total		26,2%	21,1%	25,9%	26,8%	100,0%

Chi-Square Tests

	Value	df	Asymptotic Significance (2-sided)
Pearson Chi-Square	3,635 ^a	6	,726
Likelihood Ratio	3,588	6	,732
Linear-by-Linear Association	,001	1	,982
N of Valid Cases	313		

a. 0 cells (0,0%) have expected count less than 5. The minimum expected count is 12,65.

Symmetric Measures

	Value	Approximate Significance
Nominal by Nominal	Phi	,108
	Cramer's V	,076
N of Valid Cases	313	

Kruskal-Wallis Test

Ranks

Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)		N	Mean Rank
What is your educational level?	DI palankus	82	156,53
	Žm palankus	66	159,54
	DI nepalankus	81	166,88
	Žm nepalankus	84	145,94
	Total	313	

Test Statistics^{a,b}

What is your educational level?

Kruskal-Wallis H	2,939
df	3
Asymp. Sig.	,401

a. Kruskal Wallis Test

b. Grouping Variable:

Scenarijai bendri vyrų ir moterų grupėje (4: AI/P, AI/N, Z/P, Z/N)

Descriptives

What is your age?

	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
DI palankus	82	26,7439	6,92430	,76466	25,2225	28,2653	18,00	54,00
Žm palankus	66	27,0303	7,58332	,93344	25,1661	28,8945	18,00	55,00
DI nepalankus	81	26,8765	6,56198	,72911	25,4256	28,3275	19,00	54,00
Žm nepalankus	84	26,8571	6,75410	,73693	25,3914	28,3229	18,00	54,00
Total	313	26,8690	6,90030	,39003	26,1016	27,6364	18,00	55,00

Tests of Homogeneity of Variances

		Levene Statistic	df1	df2	Sig.
What is your age?	Based on Mean	,268	3	309	,849
	Based on Median	,142	3	309	,935
	Based on Median and with adjusted df	,142	3	304,574	,935
	Based on trimmed mean	,227	3	309	,877

ANOVA

What is your age?

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	3,017	3	1,006	,021	,996
Within Groups	14852,612	309	48,067		
Total	14855,629	312			

Descriptives

What is your experience with AI-driven service (e.g. in customer service, with chatbots or virtual assistants, etc.)?

	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
DI palankus	82	2,88	1,211	,134	2,61	3,14	1	5
Žm palankus	66	3,15	1,206	,148	2,86	3,45	1	5
DI nepalankus	81	3,22	1,255	,139	2,94	3,50	1	5
Žm nepalankus	84	3,04	1,166	,127	2,78	3,29	1	5
Total	313	3,07	1,211	,068	2,93	3,20	1	5

Tests of Homogeneity of Variances

		Levene Statistic	df1	df2	Sig.
What is your experience with AI-driven service (e.g. in customer service, with chatbots or virtual assistants, etc.)?	Based on Mean	,516	3	309	,672
	Based on Median	,252	3	309	,860
	Based on Median and with adjusted df	,252	3	307,221	,860
	Based on trimmed mean	,551	3	309	,648

ANOVA

What is your experience with AI-driven service (e.g. in customer service, with chatbots or virtual assistants, etc.)?

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	5,433	3	1,811	1,238	,296
Within Groups	452,158	309	1,463		
Total	457,591	312			

Oneway

Descriptives

Atvirkščiai sukoduotas grožio paslaugų naudojimosi dažnio kintamasis, ėjęs mažėjimo tvarka

	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
DI palankus	82	2,7561	1,12832	,12460	2,5082	3,0040	1,00	5,00
Žm palankus	66	2,8030	,88090	,10843	2,5865	3,0196	1,00	5,00
DI nepalankus	81	2,8272	1,14881	,12765	2,5731	3,0812	1,00	5,00
Žm nepalankus	84	2,7857	,95780	,10450	2,5779	2,9936	1,00	5,00
Total	313	2,7923	1,03698	,05861	2,6770	2,9077	1,00	5,00

Tests of Homogeneity of Variances

		Levene Statistic	df1	df2	Sig.
Atvirkščiai sukoduotas grožio paslaugų naudojimosi dažnio kintamasis, ėjęs mažėjimo tvarka	Based on Mean	4,370	3	309	,005
	Based on Median	3,638	3	309	,013
	Based on Median and with adjusted df	3,638	3	304,005	,013
	Based on trimmed mean	4,345	3	309	,005

ANOVA

Atvirkščiai sukoduotas grožio paslaugų naudojimosi dažnio kintamasis, ėjęs mažėjimo tvarka

	Sum of Squares	df	Mean Square	F	Sig.
Between Groups	,217	3	,072	,067	,978
Within Groups	335,284	309	1,085		
Total	335,502	312			

ANOVA Effect Sizes^{a,b}

		Point Estimate	95% Confidence Interval	
			Lower	Upper
Atvirkščiai sukoduotas grožio paslaugų naudojimosi dažnio kintamasis, ėjęs mažėjimo tvarka	Eta-squared	,001	,000	,000
	Epsilon-squared	-,009	-,010	-,010
	Omega-squared Fixed-effect	-,009	-,010	-,010
	Omega-squared Random-effect	-,003	-,003	-,003

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.

Robust Tests of Equality of Means

Atvirkščiai sukoduotas grožio paslaugų naudojimosi dažnio kintamasis, ėjęs mažėjimo tvarka

	Statistic ^a	df1	df2	Sig.
Welch	,057	3	170,116	,982
Brown-Forsythe	,068	3	300,926	,977

a. Asymptotically F distributed.

14 priedas. Kolmogorovo-Smirnov testas

Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Morality total	,118	313	<,001	,933	313	<,001
Agency total	,096	313	<,001	,956	313	<,001
Fairness total	,161	313	<,001	,899	313	<,001
Satisfaction total	,216	313	<,001	,859	313	<,001

a. Lilliefors Significance Correction

15 priedas. Koreliacinė analizė

Nonparametric Correlations

Correlations			Morality total	Agency total	Fairness total	Satisfaction total
Spearman's rho	Morality total	Correlation Coefficient	1,000	-,341**	,631**	,752**
		Sig. (2-tailed)	.	<,001	<,001	<,001
		N	313	313	313	313
	Agency total	Correlation Coefficient	-,341**	1,000	-,041	-,218**
		Sig. (2-tailed)	<,001	.	,472	<,001
		N	313	313	313	313
	Fairness total	Correlation Coefficient	,631**	-,041	1,000	,717**
		Sig. (2-tailed)	<,001	,472	.	<,001
		N	313	313	313	313
	Satisfaction total	Correlation Coefficient	,752**	-,218**	,717**	1,000
		Sig. (2-tailed)	<,001	<,001	<,001	.
		N	313	313	313	313

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

16 priedas. Tarpgrupinė dispersinė analizė (ANOVA)

Univariate, two tailed

Between-Subjects Factors

		Value Label	N
Bendra grupė moterų ir vyrų (susideda iš: AI ir žmogus)	1	DI	163
	2	Žmogus	150
Bendra grupė moterų ir vyrų (susideda iš: spr. palankus vs nepalankus)	1	Palankus	148
	2	Nepalankus	165

Levene's Test of Equality of Error Variances^{a,b}

		Levene Statistic	df1	df2	Sig.
Morality total	Based on Mean	2,968	3	309	,032
	Based on Median	2,886	3	309	,036
	Based on Median and with adjusted df	2,886	3	308,095	,036
	Based on trimmed mean	3,074	3	309	,028

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Dependent variable: Morality total

b. Design: Intercept + BendDm + BENDfav + BendDm * BENDfav

Tests of Between-Subjects Effects

Dependent Variable: Morality total

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
Corrected Model	92,120 ^a	3	30,707	38,425	<,001	,272
Intercept	1773,822	1	1773,822	2219,652	<,001	,878
BendDm	3,018	1	3,018	3,776	,053	,012
BENDfav	89,895	1	89,895	112,489	<,001	,267
BendDm * BENDfav	,214	1	,214	,267	,606	,001
Error	246,935	309	,799			
Total	2082,278	313				
Corrected Total	339,056	312				

a. R Squared = ,272 (Adjusted R Squared = ,265)

Descriptive Statistics

Dependent Variable: Morality total

Bendra grupė moterų ir vyrų (susideda iš: AI ir žmogus)	Bendra grupė moterų ir vyrų (susideda iš: spr. palankus vs nepalankus)	Mean	Std. Deviation	N
DI	Palankus	2,8577	,96407	82
	Nepalankus	1,7284	,80728	81
	Total	2,2965	1,05231	163
Žmogus	Palankus	3,0025	,78146	66
	Nepalankus	1,9782	,98066	84
	Total	2,4289	1,03073	150
Total	Palankus	2,9223	,88735	148
	Nepalankus	1,8556	,90571	165
	Total	2,3600	1,04246	313

Levene's Test of Equality of Error Variances^{a,b}

		Levene Statistic	df1	df2	Sig.
Agency total	Based on Mean	,127	3	309	,944
	Based on Median	,169	3	309	,917
	Based on Median and with adjusted df	,169	3	294,098	,917
	Based on trimmed mean	,179	3	309	,910

Tests the null hypothesis that the error variance of the dependent variable is equal across groups.

a. Dependent variable: Agency total

b. Design: Intercept + BendDm + BENDfav + BendDm * BENDfav

Tests of Between-Subjects Effects

Dependent Variable: Agency total

Source	Type III Sum of Squares	df	Mean Square	F	Sig.	Partial Eta Squared
Corrected Model	44,858 ^a	3	14,953	15,960	<,001	,134
Intercept	3768,232	1	3768,232	4021,954	<,001	,929
BendDm	33,625	1	33,625	35,889	<,001	,104
BENDfav	13,238	1	13,238	14,130	<,001	,044
BendDm * BENDfav	1,574	1	1,574	1,680	,196	,005
Error	289,507	309	,937			
Total	4201,938	313				
Corrected Total	334,365	312				

a. R Squared = ,134 (Adjusted R Squared = ,126)

Descriptive Statistics

Dependent Variable: Agency total

Bendra grupė moterų ir vyrų (susideda iš: AI ir žmogus)	Bendra grupė moterų ir vyrų (susideda iš: spr. palankus vs nepalankus)	Mean	Std. Deviation	N
DI	Palankus	3,6799	,93068	82
	Nepalankus	3,9506	1,01506	81
	Total	3,8144	,97996	163
Žmogus	Palankus	2,8788	,96201	66
	Nepalankus	3,4345	,96171	84
	Total	3,1900	,99777	150
Total	Palankus	3,3226	1,02281	148
	Nepalankus	3,6879	1,01865	165
	Total	3,5152	1,03522	313

17 priedas. Tiesinės paprastosios ir tiesinės daugialypės regresijos analizės

Sąžiningumo:

Regression

Descriptive Statistics

	Mean	Std. Deviation	N
Fairness total	2,2726	1,14395	313
Morality total	2,3600	1,04246	313
Agency total	3,5152	1,03522	313

Correlations

		Fairness total	Morality total	Agency total
Pearson Correlation	Fairness total	1,000	,594	,017
	Morality total	,594	1,000	-,296
	Agency total	,017	-,296	1,000
Sig. (1-tailed)	Fairness total	.	<,001	,379
	Morality total	,000	.	,000
	Agency total	,379	,000	.
N	Fairness total	313	313	313
	Morality total	313	313	313
	Agency total	313	313	313

Variables Entered/Removed^a

Model	Variables Entered	Variables Removed	Method
1	Agency total, Morality total ^b	.	Enter

a. Dependent Variable: Fairness total

b. All requested variables entered.

Variables Entered/Removed^a

Model	Variables Entered	Variables Removed	Method
1	Agency total, Morality total ^b	.	Enter

a. Dependent Variable: Fairness total

b. All requested variables entered.

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	,628 ^a	,394	,390	,89328

a. Predictors: (Constant), Agency total, Morality total

b. Dependent Variable: Fairness total

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	160,925	2	80,462	100,836	<,001 ^b
	Residual	247,366	310	,798		
	Total	408,291	312			

a. Dependent Variable: Fairness total

b. Predictors: (Constant), Agency total, Morality total

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)	-,252	,249		-1,011	,313	-,742	,238					
	Morality total	,721	,051	,657	14,196	<,001	,621	,821	,594	,628	,628	,912	1,096
	Agency total	,234	,051	,212	4,578	<,001	,133	,335	,017	,252	,202	,912	1,096

a. Dependent Variable: Fairness total

Collinearity Diagnostics^a

Model	Dimension	Eigenvalue	Condition Index	Variance Proportions		
				(Constant)	Morality total	Agency total
1	1	2,813	1,000	,01	,02	,01
	2	,161	4,184	,01	,56	,17
	3	,026	10,372	,99	,42	,82

a. Dependent Variable: Fairness total

Casewise Diagnostics^a

Case Number	Std. Residual	Fairness total	Predicted Value	Residual
51	3,424	5,00	1,9417	3,05826
131	3,762	5,00	1,6398	3,36020
195	3,893	5,00	1,5227	3,47726
197	3,388	4,67	1,6398	3,02687

a. Dependent Variable: Fairness total

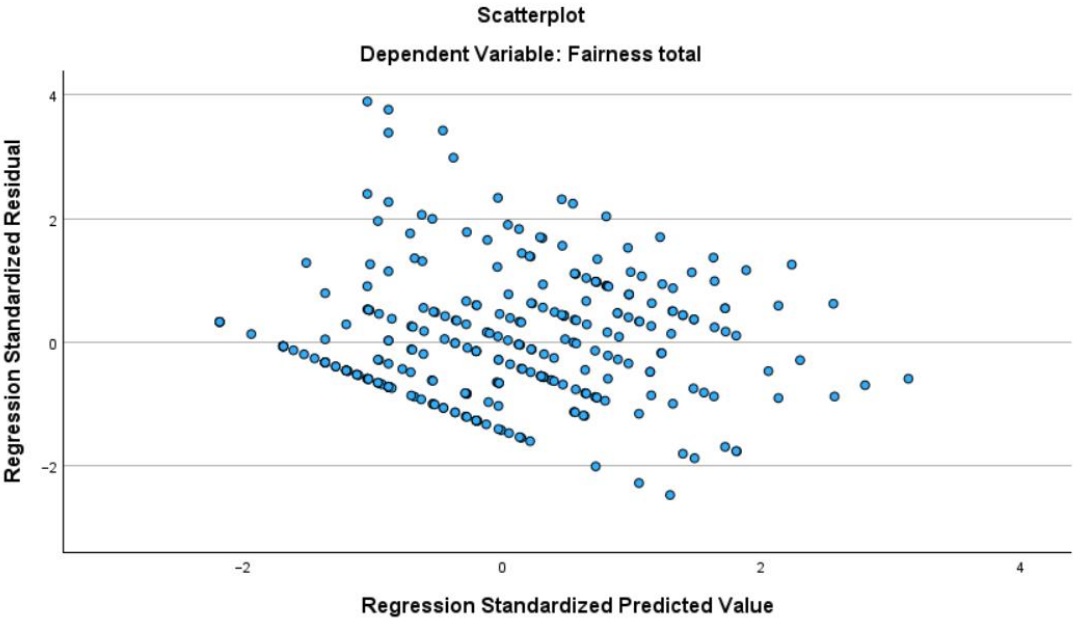
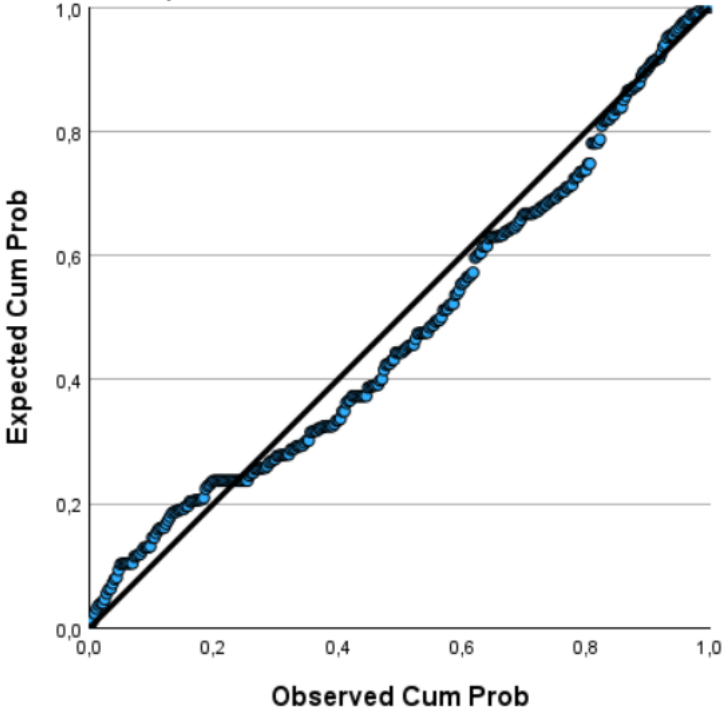
Residuals Statistics^a

	Minimum	Maximum	Mean	Std. Deviation	N
Predicted Value	,7033	4,5236	2,2726	,71818	313
Residual	-2,20185	3,47726	,00000	,89042	313
Std. Predicted Value	-2,185	3,134	,000	1,000	313
Std. Residual	-2,465	3,893	,000	,997	313

a. Dependent Variable: Fairness total

Charts

Normal P-P Plot of Regression Standardized Residual
Dependent Variable: Fairness total



Suvokiamas sąžiningumas – moralumas – sprendimo priėmėjas

Sample
Size: 313

OUTCOME VARIABLE:
Fairn

Model Summary

R	R-sq	MSE	F	df1	df2	p
,6325	,4000	,7928	68,6760	3,0000	309,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	1,9103	,3841	4,9738	,0000	1,1546	2,6660
Morality	,4473	,1506	2,9697	,0032	,1509	,7436
BendDm	-,8242	,2510	-3,2838	,0011	-1,3181	-,3303
Int_1	,1493	,0971	1,5374	,1252	-,0418	,3403

Product terms key:

Int_1 : Morality x BendDm

Test(s) of highest order unconditional interaction(s):

R2-chng	F	df1	df2	p	
X*W	,0046	2,3635	1,0000	309,0000	,1252

Focal predict: Morality (X)
Mod var: BendDm (W)

Suvokiamas sąžiningumas – veiksnas – sprendimo priėmėjas

Sample
Size: 313

OUTCOME VARIABLE:
Fairn

Model Summary

R	R-sq	MSE	F	df1	df2	p
,1710	,0292	1,2827	3,1013	3,0000	309,0000	,0270

Model

	coeff	se	t	p	LLCI	ULCI
constant	2,9844	,7797	3,8278	,0002	1,4503	4,5186
Agency	-,0311	,2040	-,1525	,8789	-,4326	,3704
BendDm	-,3865	,4737	-,8159	,4152	-1,3185	,5456
Int_1	-,0061	,1300	-,0469	,9626	-,2618	,2496

Product terms key:

Int_1 : Agency x BendDm

Test(s) of highest order unconditional interaction(s):

	R2-chng	F	df1	df2	p
X*W	,0000	,0022	1,0000	309,0000	,9626

Focal predict: Agency (X)
Mod var: BendDm (W)

Pasitenkinimo:

Regression

Descriptive Statistics

	Mean	Std. Deviation	N
Satisfaction total	2,3600	1,33747	313
Fairness total	2,2726	1,14395	313

Correlations

		Satisfaction total	Fairness total
Pearson Correlation	Satisfaction total	1,000	,669
	Fairness total	,669	1,000
Sig. (1-tailed)	Satisfaction total	.	<,001
	Fairness total	,000	.
N	Satisfaction total	313	313
	Fairness total	313	313

Variables Entered/Removed^a

Model	Variables Entered	Variables Removed	Method
1	Fairness total ^b	.	Enter

a. Dependent Variable: Satisfaction total

b. All requested variables entered.

Model Summary^b

Model	R	R Square	Adjusted R Square	Std. Error of the Estimate
1	,669 ^a	,448	,446	,99558

a. Predictors: (Constant), Fairness total

b. Dependent Variable: Satisfaction total

Collinearity Diagnostics^a

Model	Dimension	Eigenvalue	Condition Index	Variance Proportions		
				(Constant)	Morality total	Agency total
1	1	2,813	1,000	,01	,02	,01
	2	,161	4,184	,01	,56	,17
	3	,026	10,372	,99	,42	,82

a. Dependent Variable: Satisfaction total

Casewise Diagnostics^a

Case Number	Std. Residual	Satisfaction total	Predicted Value	Residual
197	4,352	5,00	1,1057	3,89425
272	-3,147	1,00	3,8158	-2,81576
274	-3,101	1,00	3,7744	-2,77437

a. Dependent Variable: Satisfaction total

Residuals Statistics^a

	Minimum	Maximum	Mean	Std. Deviation	N
Predicted Value	,9402	4,9755	2,3600	,99667	313
Residual	-2,81576	3,89425	,00000	,89189	313
Std. Predicted Value	-1,425	2,624	,000	1,000	313
Std. Residual	-3,147	4,352	,000	,997	313

a. Dependent Variable: Satisfaction total

ANOVA^a

Model		Sum of Squares	df	Mean Square	F	Sig.
1	Regression	249,854	1	249,854	252,076	<,001 ^b
	Residual	308,258	311	,991		
	Total	558,111	312			

a. Dependent Variable: Satisfaction total

b. Predictors: (Constant), Fairness total

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)	,582	,125		4,645	<,001	,336	,829					
	Fairness total	,782	,049	,669	15,877	<,001	,685	,879	,669	,669	,669	1,000	1,000

a. Dependent Variable: Satisfaction total

Collinearity Diagnostics^a

Model	Dimension	Eigenvalue	Condition Index	Variance Proportions	
				(Constant)	Fairness total
1	1	1,894	1,000	,05	,05
	2	,106	4,217	,95	,95

a. Dependent Variable: Satisfaction total

Casewise Diagnostics^a

Case Number	Std. Residual	Satisfaction total	Predicted Value	Residual
51	-3,509	1,00	4,4935	-3,49350
131	-3,509	1,00	4,4935	-3,49350
195	-3,509	1,00	4,4935	-3,49350
251	-3,247	1,00	4,2327	-3,23275
283	-3,247	1,00	4,2327	-3,23275

a. Dependent Variable: Satisfaction total

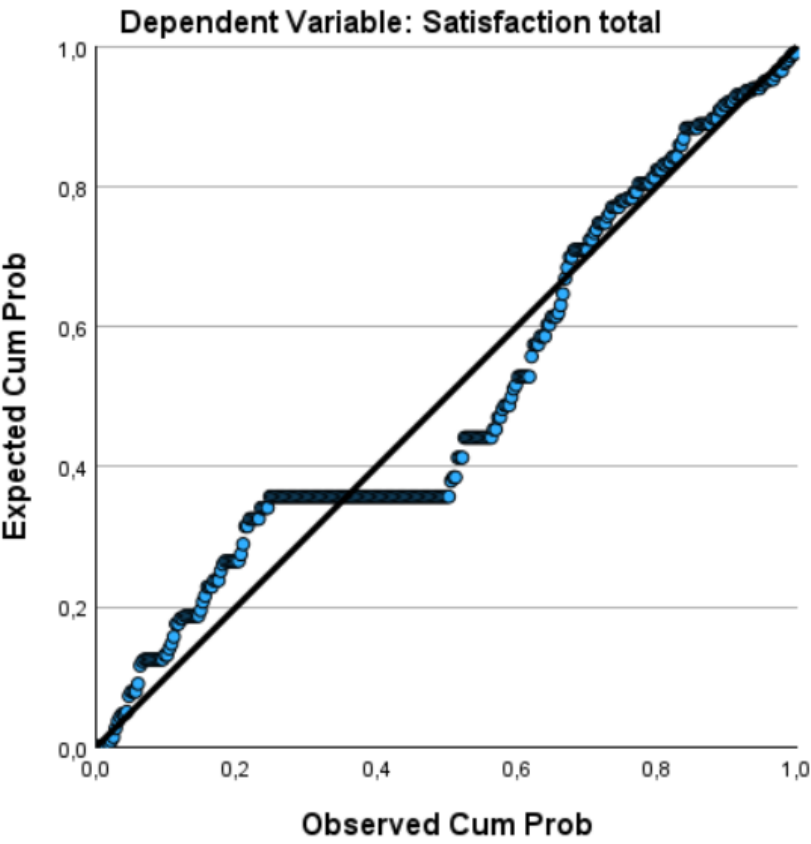
Residuals Statistics^a

	Minimum	Maximum	Mean	Std. Deviation	N
Predicted Value	1,3644	4,4935	2,3600	,89488	313
Std. Predicted Value	-1,112	2,384	,000	1,000	313
Standard Error of Predicted Value	,056	,146	,077	,019	313
Adjusted Predicted Value	1,3478	4,5699	2,3611	,89704	313
Residual	-3,49350	2,33180	,00000	,99398	313
Std. Residual	-3,509	2,342	,000	,998	313
Stud. Residual	-3,547	2,346	-,001	1,003	313
Deleted Residual	-3,56995	2,34017	-,00117	1,00271	313
Stud. Deleted Residual	-3,615	2,364	-,001	1,008	313
Mahal. Distance	,003	5,684	,997	1,057	313
Cook's Distance	,000	,138	,004	,016	313
Centered Leverage Value	,000	,018	,003	,003	313

a. Dependent Variable: Satisfaction total

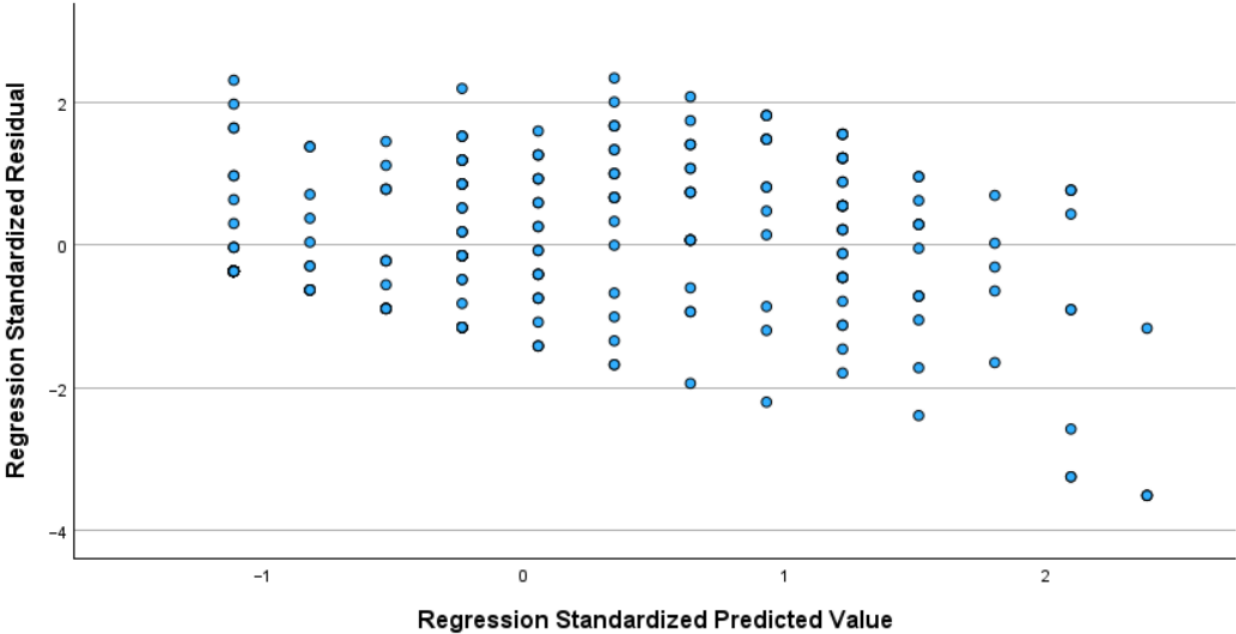
Charts

Normal P-P Plot of Regression Standardized Residual



Scatterplot

Dependent Variable: Satisfaction total



Pasitenkinimas – suvokiamas sąžiningumas – sprendimo priėmėjas

Model: 1
Y: Satis
X: Fairn
W: BendDm

Sample
Size: 313

OUTCOME VARIABLE:
Satis

Model Summary

R	R-sq	MSE	F	df1	df2	p
,6720	,4516	,9906	84,8041	3,0000	309,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	,9500	,4027	2,3590	,0189	,1576	1,7424
Fairn	,5856	,1530	3,8276	,0002	,2846	,8867
BendDm	-,2557	,2530	-1,0106	,3130	-,7534	,2421
Int_1	,1400	,1008	1,3889	,1659	-,0583	,3383

Product terms key:

Int_1 : Fairn x BendDm

Test(s) of highest order unconditional interaction(s):

	R2-chng	F	df1	df2	p
X*W	,0034	1,9291	1,0000	309,0000	,1659

Focal predict: Fairn (X)
Mod var: BendDm (W)

Mediavimo efektai. Moralumas.

Model: 4
Y: Satis
X: Morality
M: Fairn

Sample
Size: 313

OUTCOME VARIABLE:
Fairn

Model Summary

R	R-sq	MSE	F	df1	df2	p
,5943	,3532	,8492	169,8179	1,0000	311,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	,7336	,1291	5,6831	,0000	,4796	,9876
Morality	,6522	,0500	13,0314	,0000	,5537	,7506

OUTCOME VARIABLE:
Satis

Model Summary

R	R-sq	MSE	F	df1	df2	p
,7961	,6338	,6594	268,2145	2,0000	310,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	-,1949	,1195	-1,6309	,1039	-,4300	,0402
Morality	,6881	,0548	12,5500	,0000	,5803	,7960
Fairn	,4096	,0500	8,1971	,0000	,3113	,5079

***** TOTAL EFFECT MODEL *****

***** TOTAL EFFECT MODEL *****

OUTCOME VARIABLE:

Satis

Model Summary

R	R-sq	MSE	F	df1	df2	p
,7446	,5544	,7997	386,8905	1,0000	311,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	,1056	,1253	,8427	,4000	-,1409	,3520
Morality	,9553	,0486	19,6695	,0000	,8597	1,0508

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y

Effect	se	t	p	LLCI	ULCI
,9553	,0486	19,6695	,0000	,8597	1,0508

Direct effect of X on Y

Effect	se	t	p	LLCI	ULCI
,6881	,0548	12,5500	,0000	,5803	,7960

Indirect effect(s) of X on Y:

	Effect	BootSE	BootLLCI	BootULCI
Fairn	,2671	,0542	,1684	,3842

***** 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 -----

Veiksnumas.

Model: 4
Y: Satis
X: Agency
M: Fairn

Sample
Size: 313

OUTCOME VARIABLE:
Fairn

Model Summary

R	R-sq	MSE	F	df1	df2	p
,0175	,0003	1,3124	,0949	1,0000	311,0000	,7582

Model

	coeff	se	t	p	LLCI	ULCI
constant	2,2048	,2296	9,6047	,0000	1,7531	2,6564
Agency	,0193	,0627	,3081	,7582	-,1040	,1426

OUTCOME VARIABLE:
Satis

Model Summary

R	R-sq	MSE	F	df1	df2	p
,6991	,4888	,9203	148,2104	2,0000	310,0000	,0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	1,4939	,2189	6,8248	,0000	1,0632	1,9246
Agency	-,2620	,0525	-4,9940	,0000	-,3653	-,1588
Fairn	,7864	,0475	16,5614	,0000	,6930	,8798

***** TOTAL EFFECT MODEL *****

***** TOTAL EFFECT MODEL *****

OUTCOME VARIABLE:

Satis

Model Summary

R	R-sq	MSE	F	df1	df2	p
,1911	,0365	1,7290	11,7851	1,0000	311,0000	,0007

Model

	coeff	se	t	p	LLCI	ULCI
constant	3,2277	,2635	12,2504	,0000	2,7093	3,7462
Agency	-,2469	,0719	-3,4329	,0007	-,3884	-,1054

***** TOTAL, DIRECT, AND INDIRECT EFFECTS OF X ON Y *****

Total effect of X on Y

Effect	se	t	p	LLCI	ULCI
-,2469	,0719	-3,4329	,0007	-,3884	-,1054

Direct effect of X on Y

Effect	se	t	p	LLCI	ULCI
-,2620	,0525	-4,9940	,0000	-,3653	-,1588

Indirect effect(s) of X on Y:

Effect	BootSE	BootLLCI	BootULCI
Fairn	,0152	,0520	-,0873

***** 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 -----