

**Vilniaus universiteto Teisės fakulteto
Privatinės teisės katedra**

Gabijos Leckaitės,
V kurso, darbo ir socialinės teisės studijų
šakos studentės

Magistro darbas

Dirbtinis intelektas ir naujosios technologijos: darbo teisės transformacija
Artificial intelligence and emerging technologies: the transformation of labor law

Vadovas: doc. dr. Tomas Bagdanskis
Recenzentas: doc. dr. Justinas Usonis

Vilnius
2025

ANOTACIJA IR PAGRINDINIAI ŽODŽIAI

Šiame darbe analizuojamas dirbtinio intelekto technologijų poveikis darbo teisės santykiams ir teisiniam reguliavimui. Tyrime nagrinėjamas algoritminio valdymo reiškiny, kai darbo procesai organizuojami, prižiūrimi ir vertinami automatizuotomis priemonėmis, grindžiamomis duomenų analize ir dirbtinio intelekto sprendimais. Atskleidžiami esminiai darbuotojų teisių iššūkiai, kylantys dėl sprendimų skaidrumo, duomenų apsaugos ir teisės būti nediskriminuojamam. Analizuojama, kaip Bendrojo duomenų apsaugos reglamento 22 straipsnis taikomas darbo santykiuose, ir kokias garantijas darbuotojui suteikia. Įvertinamas Europos Sąjungos dirbtinio intelekto akto rizikos pagrindu grindžiamas požiūris ir jo poveikis darbuotojų teisių apsaugos sistemai.

Pagrindiniai žodžiai: algoritminis valdymas, dirbtinis intelektas, darbo teisė, duomenų apsauga, teisinis reguliavimas

This thesis examines the impact of artificial intelligence technologies on labour law relations and legal regulation. The study focuses on the phenomenon of algorithmic management, where work processes are organised, monitored, and evaluated using automated tools based on data analytics and AI-driven decision-making. It identifies key challenges to employees' rights arising from the lack of transparency in decisions, data protection issues, and the right not to be discriminated against. The research analyses the application of Article 22 of the General Data Protection Regulation (GDPR) in employment relations and the guarantees it provides to employees. It also evaluates the risk-based approach enshrined in the European Union's Artificial Intelligence Act and its implications for the employee protection framework.

Keywords: algorithmic management, artificial intelligence, labour law, data protection, legal regulation

SANTRUMPOS

Asmens duomenys - bet kokia informacija, susijusi su identifiukuotu ar galimu identifiukuoti fiziniu asmeniu, kaip tai apibrėžta Bendrajame duomenų apsaugos reglamente (BDAR).

Baltoji knyga dėl dirbtinio intelekto – 2020 m. vasario 19 d. Europos Komisijos paskelbtas politinės strategijos dokumentas *White Paper on Artificial Intelligence – A European Approach to Excellence and Trust*, kuriuo siekta apibrėžti kryptį, kaip ES galėtų skatinti inovacijas DI srityje ir kartu užtikrinti aukštą žmogaus teisių, saugumo bei duomenų apsaugos standartą. Dokumente buvo siūloma kurti rizikomis grįstą reguliavimo modelį, ypatingą dėmesį skiriant didelės rizikos DI sistemoms (Europos Komisija, 2020).

BDAR - Bendrasis duomenų apsaugos reglamentas (Reglamentas (ES) 2016/679) – pagrindinis ES teisės aktas, reglamentuojantis fizinių asmenų duomenų tvarkymą bei nustatantis jų apsaugos standartus visose ES valstybėse narėse.

DK - Lietuvos Respublikos darbo kodeksas – pagrindinis Lietuvos teisės aktas, reglamentuojantis darbo santykius, darbuotojų teises, pareigas ir socialines garantijas.

DI - Dirbtinis intelektas – pažangių technologijų sistema, gebanti atlikti funkcijas, kurios įprastai reikalauja žmogaus intelekto, tokias kaip mokymasis, sprendimų priėmimas, kalbos atpažinimas ar prognozavimas.

DI aktas - Eur Europos Parlamento ir Tarybos reglamentas (ES) 2024/1689, pirmasis išsamus ES teisės aktas, reglamentuojantis dirbtinio intelekto sistemų vystymą, tiekimą rinkai ir naudojimą, siekiant užtikrinti pagrindinių teisių, saugumo ir inovacijų pusiausvyrą.

Dirbtinio intelekto atsakomybės direktyvos pasiūlymas /AILD - Dirbtinio intelekto atsakomybės direktyvos pasiūlymas (Artificial Intelligence Liability Directive) – Europos Komisijos 2022 m. pateiktas teisėkūros pasiūlymas, siekiantis suderinti griežtosios atsakomybės principus dėl žalos, padarytos naudojant DI sistemas.

EK Aukšto lygio ekspertų grupė dėl dirbtinio intelekto (AI HLEG) – Europos Komisijos įsteigta nepriklausoma ekspertų grupė, veikusi 2018–2020 m., kurios pagrindinis tikslas buvo teikti rekomendacijas dėl patikimo, į žmogų orientuoto dirbtinio intelekto kūrimo ir diegimo Europos Sąjungoje.

Griežtoji atsakomybė - civilinės atsakomybės rūšis, kuriai taikyti nereikalaujama įrodyti žalą padariusio asmens kaltės. Ši atsakomybės forma grindžiama veiklos pavojingumu, techniniu sudėtingumu arba socialine rizika ir taikoma nepriklausomai nuo to, ar žalą padaręs subjektas elgėsi neteisėtai ar nerūpestingai.

Platforminio darbo direktyvos projektas - 2021 m. gruodžio mėn. Europos Komisijos pasiūlyta direktyva dėl dirbančiųjų per skaitmenines darbo platformas teisinio statuso, teisių apsaugos ir atsakomybės paskirstymo.

TDO / ILO (Tarptautinė darbo organizacija / International Labour Organization) - Jungtinių Tautų specializuota agentūra, įkurta 1919 m., atsakinga už tarptautinių darbo normų kūrimą, stebėseną ir socialinės apsaugos vystymą pasauliniu mastu. TDO yra parengusi daug svarbių dokumentų, tarp jų – Konvenciją Nr. 102 dėl socialinės apsaugos minimalių standartų, taip pat R202 rekomendaciją dėl socialinės apsaugos pagrindo (ILO, 2012). Pastaraisiais metais

ILO taip pat analizuoja DI įtaką darbo rinkai, pabrėždama algoritminio valdymo ir automatizacijos poveikį darbuotojų teisėms ir socialinėms garantijoms.

Vilniaus konvencija - Europos Tarybos parengta Dirbtinio intelekto ir žmogaus teisių, demokratijos ir teisinės valstybės pagrindų konvencija (dar vadinama Dirbtinio intelekto konvencija arba Vilniaus konvencija) pasirašyta 2024 m. rugsėjo 5 d. Vilniuje.

TURINYS

SANTRUMPOS	3
IŽANGA.....	6
1. DIRBTINIO INTELEKTO TEISINĖ SAMPRATA IR INTEGRACIJA Į DARBO SANTYKIUS	8
1.1. Dirbtinio intelekto sąvoka, veikimo principai ir reguliavimo struktūra	8
1.2. Dirbtinio intelekto taikymas darbo teisiniuose santykiuose	11
1.3. Lietuvos Respublikos darbo kodeksas ir jo santykis su DI reguliavimu	17
2. DIRBTINIO INTELEKTO POVEIKIS DARBUOTOJŲ TEISĖMS	26
2.1. Biometrinių duomenų naudojimas ir privatumo teisės apsauga	26
2.2. Darbuotojų stebėseną, vertinimas, atrankų procesas	35
2.3. Automatizuotų sprendimų poveikis	44
3. ATSAKOMYBĖS PASKIRSTYMAS UŽ DIRBTINIO INTELEKTO NAUDOJIMĄ DARBO TEISINIUOSE SANTYKIUOSE	50
3.1. Darbdavių, darbuotojų ir trečiųjų šalių atsakomybės aspektai	50
3.2. Civilinės atsakomybės už DI klaidas ir diskriminaciją modeliai	57
3.3. Vilniaus konvencijos poveikis atsakomybės teisiniam reguliavimui	65
IŠVADOS	72
ŠALTINIŲ SĄRAŠAS	74
SANTRAUKA	82
SUMMARY	83

IŽANGA

Temos aktualumas. Dirbtinio intelekto (DI) ir naujų technologijų plėtra keičia tradicinius darbo santykių modelius ir kelia esminius iššūkius tiek darbdaviams, tiek darbuotojams, tiek darbuotojų teisių apsaugai. DI sprendimai, automatizuotas darbuotojų atrankos procesas, veiklos stebėjimas bei rezultatų analizė darbe – visa tai tampa kasdienybe šiuolaikinėse organizacijose. Šie pokyčiai kelia ne tik praktinius, bet ir teisinius klausimus, susijusius su darbuotojų teisėmis, jų privatumu, atsakomybe bei darbo sąlygų teisingumu. Lietuvoje teisės aktai dar nėra iki galo prisitaikę prie šių naujų iššūkių, todėl aktualu analizuoti, ar esama teisinė bazė užtikrina pakankamą darbuotojų apsaugą ir teisinį aiškumą taikant DI sprendimus. Tiek nacionaliniu, tiek Europos Sąjungos lygiu vis dar diskutuojama dėl efektyviausių reguliavimo formų, o mokslinėje literatūroje šie klausimai dažnai nagrinėjami fragmentiškai.

Darbo tikslas. Įvertinti esamą dirbtinio intelekto ir naujų technologijų teisinį reguliavimą darbo teisės srityje, kartu analizuojant, ar egzistuojančios normos sudaro pakankamą pagrindą darbuotojų teisių apsaugai ir skaidriam sprendimų priėmimui, ar būtinas naujas teisinis modelis bei išanalizuoti papildomo ar naujo reglamentavimo poreikį Lietuvoje.

Darbo uždaviniai.

1. Apibrėžti dirbtinio intelekto sampratą ir išanalizuoti jo taikymo sritis darbo teisėje;
2. Išanalizuoti DI technologijų poveikį darbuotojų teisėms (privatumui, atrankai, vertinimui);
3. Įvertinti Lietuvos Respublikos darbo kodekso ir BDAR normų taikymo galimybes DI kontekste;
4. Įvertinti ES teisės, DI akto įtaką nacionaliniam reguliavimui;
5. Nustatyti esamas teises spragas bei atsakomybės taikymo problemas dėl DI naudojimo;
6. Pateikti siūlymus dėl nacionalinio teisinio reguliavimo tobulinimo.

Objektas. Darbo objektas – dirbtinio intelekto ir susijusių technologijų taikymo darbo santykiuose teisinis reguliavimas. Dėmesys sutelkiamas į nacionalinės teisės normų analizę, taip pat vertinama Europos Sąjungos iniciatyvų, ypač Dirbtinio intelekto akto, įtaka. Kiti DI poveikio aspektai, tokie kaip civilinė atsakomybė ar intelektinė nuosavybė, analizuojami tiek, kiek jie susiję su darbo teisės problematika.

Tyrimo metodai. Tyrime taikomas sisteminis, lyginamasis, loginis ir teisės doktrinos analizės metodai. Sisteminis metodas leidžia atskleisti normų tarpusavio sąveiką, lyginamasis –

identifikuoti tarptautines teisinės tendencijas, o doktrinos analizė – pagrįsti moksliniais šaltiniais galimus reguliacinės sistemos pokyčius.

Darbo originalumas. Lietuvoje DI taikymo darbo teisėje tema dar nėra išsamiai nagrinėta. Didžioji dalis esamų darbų apsiriboja bendresnėmis skaitmenizacijos ar technologinės pažangos įtakos temomis. Šis darbas išsiskiria tuo, kad koncentruojamasi į specifinį, bet vis svarbesnį kontekstą – DI sprendimų poveikį darbuotojų teisėms, jų duomenų apsaugai, automatizuotų sprendimų priėmimui bei teisinio aiškumo užtikrinimui. Taip pat darbe atsižvelgiama į naujausias ES teisėkūros tendencijas, kurios kol kas tik formuojasi, tačiau ateityje turės esminį poveikį nacionaliniam reguliavimui.

Svarbiausi šaltiniai. Darbe analizuojami Lietuvos Respublikos darbo kodeksas, Bendrasis duomenų apsaugos reglamentas (BDAR), Europos Sąjungos Dirbtinio intelekto akto projektas bei kiti nacionaliniai ir ES teisės šaltiniai. Teoriniu pagrindu pasitelkiami tiek Lietuvos, tiek užsienio autorių moksliniai darbai dirbtinio intelekto, darbo teisės, duomenų apsaugos ir skaitmeninės transformacijos srityse. Ypatingas dėmesys skiriamas aktualiai teismų praktikai.

1. DIRBTINIO INTELEKTO TEISINĖ SAMPRATA IR INTEGRACIJA Į DARBO SANTYKIUS

1.1. Dirbtinio intelekto sąvoka, veikimo principai ir reguliavimo struktūra

Dirbtinio intelekto (toliau – DI) samprata pastaraisiais dešimtmečiais tapo tiek techninių, tiek filosofinių, tiek ir teisinių debatų ašimi. Nepaisant sparčios technologinės pažangos, pati sąvoka lieka diskusijų objektu, ypač siekiant ją perkelti į teisės normų lauką. Terminas „dirbtinis intelektas“ pirmą kartą moksliskai pavartotas 1956 m. Dartmuto konferencijoje, kurios metu tyrėjai – John McCarthy, Marvin Minsky, Nathaniel Rochester ir Claude Shannon – pasiūlė studijuoti galimybę sukurti sistemas, gebančias imituoti žmogaus intelektą, mokytis ir spręsti problemas be žmogaus įsikišimo (McCarthy et al., 1955).

XXI a. pradžioje prasidėjo vadinamasis mašininio mokymosi (*machine learning*) bumai, kuris tapo pagrindiniu šiuolaikinio DI veikimo principu. Šios sistemos geba mokytis iš savo patirties, apdoroti ir analizuoti milžiniškus duomenų kiekius, todėl jos efektyviai pritaikomos įvairiose srityse – nuo finansų ir sveikatos apsaugos iki teisinės praktikos ir viešojo administravimo (Jordan ir Mitchell, 2015). XXI a. pradžioje išryškėjo mašininio mokymosi proveržis – metodas, leidžiantis sistemoms mokytis iš patirties, analizuoti duomenis ir tobulėti be tiesioginės žmogaus intervencijos. Ypatinę reikšmę įgavo gilieji neuroniniai tinklai (*deep neural networks*), struktūriškai imituojantys žmogaus smegenų veikimą. Jie apdoroja sudėtingus duomenų rinkinius (pvz., vaizdus, kalbą, elgesio modelius) per daugiapakopius sluoksnius ir sugeba suprasti duomenų struktūrą bei prasmę – taip leidžiant DI sistemoms veikti ne tik skaičiavimo, bet ir interpretavimo lygmenyje (Goodfellow et al., 2016).

Dirbtinis intelektas yra viena svarbiausių ir sparčiausiai besivystančių technologijų XXI amžiuje. Jos poveikis apima įvairias žmogaus veiklos sritis, nuo įdarbinimo iki kreditų suteikimo ar net teisėsaugos sprendimų (Tamošaitis, M. ir Gaubienė, N., 2025). Šiandien, kai vis dažniau kalbama apie dirbtinį intelektą ir jo potencialą, iškyla poreikis aiškiai apibrėžti šią sąvoką, suprasti jos veikimo principus ir suvokti galimus pavojus, susijusius su šių technologijų naudojimu. Toks poreikis kyla todėl, kad dirbtinio intelekto sistemos vis dažniau tampa autonomiškais sprendimų priėmimo priemonėmis, darančiomis teisinius ar faktinius padarinius žmogui. Neapibrėžtumai dėl DI veikimo kelia grėsmę ir darbo teisiniuose santykiuose dėl sprendimų aiškumo, atsakomybės dėl DI naudojimo ir tokių sprendimų apskundimo galimybės.

Valstybės duomenų agentūra DI apibrėžia kaip sistemas, naudojančias tokias technologijas kaip teksto gavyba, kompiuterinė rega, kalbos atpažinimas, natūraliosios kalbos generavimas, gilusis mokymasis ir kt., kurios renka ar naudoja duomenis tam, kad su

autonomijos elementais numatytų, nuspręstų ar rekomenduotų veiksmus konkrečioms tikslams pasiekti (Valstybės duomenų agentūra, 2025).

Vienas iš plačiausiai naudojamų yra Europos Parlamento pateiktas apibrėžimas, kuriame teigiama, kad „dirbtinis intelektas yra mašinos gebėjimas, kuris yra panašus į žmonių galimybes, tokias kaip argumentavimas, mokymasis, planavimas ir kūrybiškumas“ (Europos Parlamentas, 2020). Šis apibrėžimas pabrėžia, kad DI gali atlikti užduotis, kurias iki šiol laikėme būdingomis tik žmogui – mąstyti, mokytis ir spręsti problemas.

Tuo tarpu Kaplanas ir Haenleinas (2019) DI apibrėžia kaip sistemą, kuri sugeba suprasti informaciją iš aplinkos, mokytis iš patirties ir šias žinias pritaikyti, kad pasiektų konkrečius tikslus. Šis požiūris labiau pabrėžia DI lankstumą ir gebėjimą prisitaikyti prie besikeičiančių aplinkybių.

Lietuvos Visuotinė lietuvių enciklopedija apibrėžia DI kaip technologijų ir metodų rinkinį, skirtą išmaniems įrenginiams ar sistemoms kurti. Šis apibrėžimas labiau techninis – jis išskiria DI kaip mokslo ir inžinerijos sritį. Tai parodo, kad šiai dienai nėra universalaus ir vienareikšmio DI apibrėžimo – skirtingi kontekstai lemia skirtingas akcentuotas savybes.

Plačiai naudojama Europos Komisijos Aukšto lygio ekspertų grupės dirbtinio intelekto klausimais (AI HLEG) definicija teigia, kad DI yra „žmonių sukurtos sistemos, kurios geba savarankiškai analizuoti informaciją, daryti išvadas, mokytis iš savo veiklos rezultatų ir prisitaikyti prie naujų situacijų“ (European Commission, 2019, p. 6). Tokia koncepcija parodo ne tik technologinį pažangumą, bet ir potencialą veikti kaip autonominiai subjektai, galintys imituoti sprendimų priėmimo, interpretavimo ir adaptacijos gebėjimus.

Teisiniu požiūriu vienas reikšmingiausių bandymų apibrėžti DI sampratą yra Europos Parlamento ir Tarybos 2024 m. AI aktas. Šio akto 3 straipsnio 1 dalyje nustatyta, kad DI sistema – tai programinė įranga, galinti gauti duomenis, juos apdoroti ir savarankiškai generuoti rezultatus, darančius įtaką žmonių elgesiui ar sprendimams (Europos Parlamentas ir Taryba, 2024). Šis apibrėžimas išskiria keturis pagrindinius požymius: autonomiją, prisitaikymą, duomenų analizę ir rezultatų generavimą. AI akto apibrėžimas yra laikytinas teisiškai reikšmingiausiu, nes jis yra vienintelis privalomas, įpareigojantis visoms ES valstybėms narėms. Kiti apibrėžimai – kaip Europos Parlamento ar Kaplan ir Haenlein – turi teorinę ar praktinę reikšmę, tačiau nėra taikytini teisiniame reguliavime. Nacionaliniu lygiu būtų tikslinga remtis šiuo apibrėžimu ir, jei Lietuvos darbo kodeksas (toliau – DK) būtų papildytas DI reglamentavimu, būtų tikslinga daryti nuorodą į DI akto sąvoką, siekiant vientisumo ir aiškumo.

DI kelia esminių etinių ir teisinių klausimų, susijusių su sprendimų skaidrumu, atsakomybe, diskriminacijos prevencija bei duomenų apsauga. Boddington (2017) ir Floridi bei Taddeo (2016) pažymi, kad svarstyтина ne tik DI priimtų sprendimų techninė pagrįstumo analizė,

bet ir jų moralinis statusas bei galimybė priskirti teisinę atsakomybę už DI veiksmus. Algoritmų šališkumas, atsirandantis dėl iškreiptų duomenų ar mokymo modelių, gali lemti diskriminuojančius ar neteisingus sprendimus (Wachter et al., 2017).

Reaguodama į šias rizikas, Europos Sąjunga 2024 m. priėmė Dirbtinio intelekto aktą, kuris remiasi rizika grįstu požiūriu. Pagal jį, DI sistemos klasifikuojamos atsižvelgiant į jų poveikį pagrindinėms teisėms ir saugumui. Ypatinę vietą užima didelės rizikos DI sistemos – tai tokios sistemos, kurių taikymas gali reikšmingai paveikti asmens pagrindines teises, sukelti teisinius padarinius ar turėti esminį poveikį asmens gyvenimo aplinkybėms.

Pagal AI akto 1 priedo 4(a) punktą, DI sistemos, taikomos darbo santykiuose, įskaitant atranką, vertinimą, užduočių paskirstymą, skatinimą ar atleidimą, laikomos didelės rizikos sistemomis. Joms taikomi papildomi reikalavimai: rizikos vertinimas, žmogaus priežiūra, sistemos dokumentacija ir sprendimų paaiškinamumas. AI akto 29 straipsnis papildomai įtvirtina, kad naudotojai (darbdaviai) turi užtikrinti darbuotojų informavimą apie DI veikimą, sprendimo logikos supratimą ir galimybę įsitraukti į sprendimo priėmimo procesą.

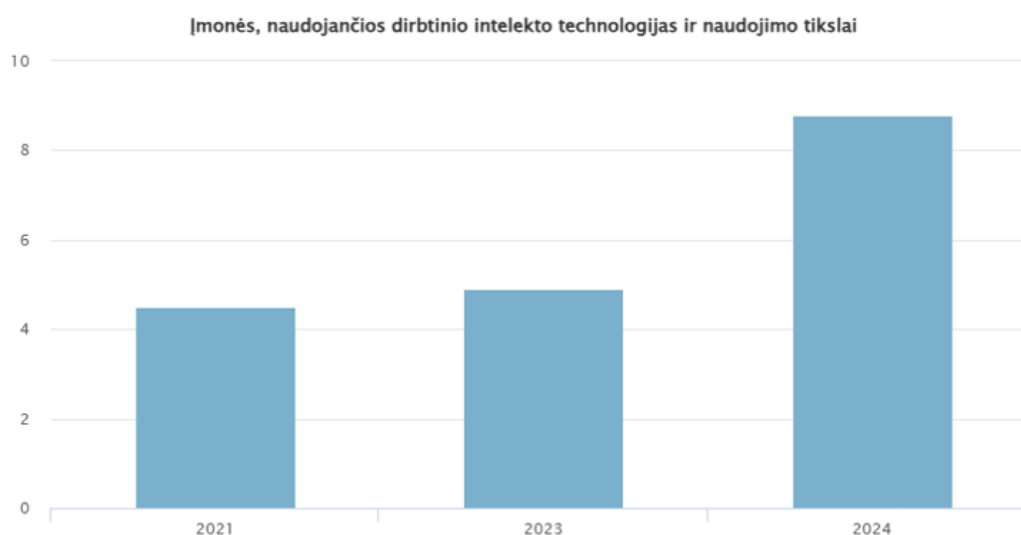
Ši DI rizikos klasifikacija keičia patį darbdavio vaidmenį – nuo technologijų naudotojo iki atsakingo subjekto, kuriam taikomi teisiniai įpareigojimai, panašūs į priežiūros institucijų reikalavimus. Taip DI reguliavimas tiesiogiai siejasi su skaidrumu ir informavimu, numatytu BDAR 13–15 straipsniuose bei AI akto 29 straipsnyje, teise į konsultacijas (tam tikrai atvejais dėl DI taikymo ir naudojimo darbovietėse), kuri įtvirtinta DK 206 straipsnyje, bei proporcingumo ir sąžiningumo principai, kylantys iš DK 24 straipsnio. Šių normų visuma sudaro pagrindą reikalauti, kad automatizuoti sprendimai būtų ne tik techniškai veiksmingi, bet ir teisiškai paaiškinami, etiški ir ginčijami.

Todėl šiuolaikinis DI reguliavimas neapsiriboja technologinėmis gairėmis, jis tampa integralia žmogaus teisių apsaugos darbo santykiuose dalimi. Dirbtinio intelekto aktas formuoja ne tik teisinę sąvoką, bet ir funkcinį veiklos modelį, kuriuo vadovaujantis DI turi būti taikomas etiškai, skaidriai ir teisėtai. Šio modelio taikymas darbo santykiuose kelia svarbių klausimų, kurie analizuojami kituose darbo skyriuose.

1.2. Dirbtinio intelekto taikymas darbo teisiniuose santykiuose

DI technologijų diegimas į darbo santykius lemia esminius pokyčius tradicinėse personalo valdymo praktikose. Šiuolaikinėje praktikoje DI vis dažniau taikomas ankstyvose atrankos stadijose, kai automatizuotos sistemos analizuoja gyvenimo aprašymus, motyvacinius laiškus, o tam tikrais atvejais ir kandidato elgseną nuotolinio vaizdo interviu metu. Tokios sistemos dažnai veikia mašininio mokymosi algoritmais, kurie treniruojami istoriniais atrankos duomenimis, taip imituodami žmogaus sprendimų logiką. Nors dirbtinio intelekto technologijos žada efektyvesnius sprendimus, mokslinėje literatūroje vis dažniau akcentuojamos rizikos, susijusios su tokių sistemų skaidrumu, interpretavimo galimybėmis ir galimu šališkumu. Ypač jautri situacija susidaro tais atvejais, kai sprendimai priimami automatiškai, be žmogaus įsikišimo, o sprendimo logika lieka nepaaiškinama arba neprieinama. Europos Komisijos Aukšto lygio ekspertų grupė savo etikos gairėse pabrėžia būtinybę užtikrinti, kad DI sprendimai būtų atsekami, paaiškinami ir nešališki, o jų poveikis žmonėms būtų aiškiai suprantamas bei suderinamas su žmogaus teisėmis ir teisinėmis garantijomis (Europos Komisija, 2019). DI sistemų integravimas į darbo santykių sritį keičia tradicinį požiūrį, kai sprendimus priima žmogus, o ne algoritmas, kurie susiję su darbuotojo atranka, veiklos stebėseną bei sprendimais dėl darbo santykių nutraukimo.

Valstybės duomenų agentūros oficialios statistikos portalo puslapyje galime matyti oficialią statistiką, kurioje pateikiami duomenys apie tai, kiek procentine išraiška visų įmonių naudoja dirbtinį intelektą darbovietėse ir kokiems tikslams šis įrankis naudojamas.



Šaltinis: Valstybės duomenų agentūra (2023), *Skaitmeninė aplinka: dirbtinis intelektas*.

Šioje lentelėje atvaizduota procentinė išraiška visų įmonių naudojančių dirbtinio intelekto technologijas, atskirai neskaidant įmonių pagal darbuotojų skaičių jose. Kaip matyti lentelėje, DI

technologijų naudojimas įmonėse 2021 metais buvo 4,5 procentai iš visų įmonių. 2023 metais šis procentas šiek tiek paaugo iki 4,9 procentų. Didžiausias augimas matomas 2024 metais, kuomet šis procentas išaugo iki 8,8 procentų. Tai suponuoja didžiulį augimą per metus laiko. DI naudojimas įmonėse sparčiai didėja, todėl galima prognozuoti, kad DI panaudojimas darbovietėse metams bėgant tik dar sparčiau didės. Tokios prognozės tik dar labiau sustiprina dirbtinio intelekto reglamentavimo poreikį darbo teisės aktuose.

Iš detalesnės duomenų statistikos lentelės, kuri pateikta 1 priede matyti, kad daugiausiai dirbtinio intelekto įrankiais naudojasi didelės įmonės, turinčios 250 ir daugiau darbuotojų. Dažniausiai taikomos DI technologijos susijusios su rašytinės kalbos analize ir įvairių darbo procesų automatizavimu. Tai rodo, kad didelės įmonės, susiduriančios su kompleksiniais darbo organizavimo iššūkiais, siekia didinti efektyvumą per automatizuotus sprendimus. Tokios tendencijos leidžia daryti išvadą, kad dirbtinis intelektas tampa nebe papildoma technologija, o esminiu darbo organizavimo įrankiu. Todėl kyla poreikis ne tik adaptuoti esamas darbo teisės nuostatas, bet ir iš esmės perkonstruoti normatyvinį pagrindą, siekiant užtikrinti, kad darbo teisė apimtų ir algoritminio valdymo procesus. Tai ypač svarbu dėl to, kad šiuo metu Lietuvos Respublikos darbo kodekse nėra specialių nuostatų, reguliuojančių automatizuotų sprendimų taikymą darbo aplinkoje. Internetu, įmonių, kurios vykdo darbuotojų atrankas ir šias paslaugas siūlo kitoms įmonėms, puslapiuose, galima pastebėti, kad personalo atrankų specialistai efektyviai naudoja DI ir siūlo DI pagrįstus sprendimus kitoms įmonėms. Dirbtinis intelektas taikomas įvairiose srityse, įskaitant logistiką, finansus, sveikatos apsaugą, mažmeninę prekybą, švietimą ir viešąjį valdymą. Logistikos sektoriuje DI padeda optimizuoti maršrutus ir tiekimo grandines; finansuose – aptikti sukčiavimą ar automatizuoti analizę; sveikatos apsaugoje – diagnozuoti ligas ir personalizuoti gydymą; o švietime – teikti individualizuotas mokymosi rekomendacijas. Kaip pažymi Europos Parlamentas (2020), DI geba pagerinti tiek viešąsias paslaugas, tiek verslo efektyvumą, tačiau tuo pat metu kelia svarbių reguliavimo klausimų, susijusių su saugumu ir atsakomybe. Tokie sprendimai pažada objektyvumą, greitį ir efektyvumą. Sutiktina, kad automatizuoti sprendimų mechanizmai, pasitelkdamie mašininį mokymąsi ir duomenų analitiką, gali didinti efektyvumą atrankų procese, bet kas nepamirima, kad toks atrankų procesas gali kelti ir naujus iššūkius, ypač susijusius su darbuotojų teisėmis, skaidrumu ir diskriminacijos prevencija.

ES DI akte (Reglamentas (ES) 2024/1689) darbuotojų veiklos vertinimas priskiriamas prie didelės rizikos sričių: tai apima tiek veiklos vertinimą, tiek elgsenos analizę, jei jos daro įtaką darbo teisiniam statusui. 1 priedo 4(a) punkte nurodoma, kad tokios sistemos turi būti skaidrios, stebimos ir pagrįstos žmogaus įsikišimu (Europos Parlamentas ir Taryba, 2024).

Galiausiai, Europos Komisijos Aukšto lygio ekspertų grupės DI gairėse (2019) rekomenduojama diegiant veiklos vertinimo DI sistemas vadovautis „etiško projektavimo“ (angl. *ethics by design*) principu – tai reiškia, kad darbuotojo autonomija, teisė į privatumą turi būti sistemingai įvertinami ir įtraukiami į sprendimų projektavimą, atsižvelgiant į galimą poveikį žmogaus teisėms. Nors gairėse nenurodoma konkreti techninė forma, kaip šie principai turi būti integruoti, rekomenduojama pasitelkti poveikio vertinimus, dalyvavimu grįstą projektavimą ir skaidrumą, mechanizmus, leidžiančius darbuotojams suprasti ir kontroliuoti DI sprendimų poveikį jų darbo situacijai (Europos Komisija, 2019).

Vienas kertinių Dirbtinio intelekto akto elementų yra rizika grįstas požiūris, pagal kurį DI sistemos skirstomos į keturias kategorijas: minimalios, ribotos, didelės ir nepriimtinos rizikos. Pastaroji kategorija yra pati griežčiausia, kadangi apima tokias praktikas, kurios laikomos nesuderinamomis su pagrindinėmis žmogaus teisėmis, orumu ir Europos Sąjungos vertybėmis, todėl yra visiškai draudžiamos visose ES valstybėse narėse (Europos Parlamentas ir Taryba, 2024, 5 str.).

2025 m. vasario 2 d. Europos Komisija paskelbė oficialias gaires dėl nepriimtinos rizikos DI praktikų, kurios padeda detaliau suprasti draudimų turinį ir taikymo ribas (Europos Komisija, 2025a). Gairėse akcentuojama, kad pagrindinis tikslas – užtikrinti, jog DI sistemų naudojimas nepažeistų žmogaus autonomijos, privatumo, lygybės bei nediskriminavimo principų. Jos skirtos tiek nacionalinėms priežiūros institucijoms, tiek DI naudotojams (darbuotojams/darbdaviams) ir kūrėjams. Nepriimtina rizika apima konkrečius, griežtai apibrėžtus DI taikymo atvejus. Tarp jų – emocijų atpažinimo technologijų taikymas darbo vietose, išskyrus atvejus, kai tai yra būtina dėl saugos priežasčių (pvz., stebint piloto nuovargį). Gairėse aiškiai pabrėžiama, kad tokios sistemos neturi būti naudojamos vertinant darbuotojo nuotaiką, produktyvumą ar tinkamumą dirbti, nes tai pažeidžia asmens psichologinį integralumą ir darbo teisės principą nediskriminuoti (Europos Komisija, 2025a).

Taip pat uždraudžiamas biometrinių duomenų naudojimas asmenų skirstymui pagal jautrias charakteristikas, tokias kaip rasė, religija, seksualinė orientacija ar politinės pažiūros. Toks draudimas ypač svarbus darbo santykiuose, kur DI sistemos gali būti naudojamos atrankai ar vertinimui – gairėse nurodoma, kad bet koks tokio pobūdžio skirstymas kelia neleistinos diskriminacijos riziką (Europos Komisija, 2025a). Kadangi DI naudojamas atrankų procese, darbdaviai privalo ne tik atsakingai pasirinkti DI sistemas, bet ir užtikrinti, kad pasitelkiami

įrankiai atitiktų duomenų apsaugos bei lygių galimybių principus. Toks naudojamas įrankis veiktų tinkamai, neskirstant asmenų pagal jautrias charakteristikas, kad įrankis nediskriminuotų atrankų kandidatų, kandidatas turėtų galimybę peržiūrėti sprendimų logiką bei būtų užtikrinamas žmogaus įsikišimas. Kitaip tokie veiksmai galėtų pažeisti Lygių galimybių įstatymo 6 straipsnio 1 dalį, kuri draudžia diskriminaciją įdarbinimo procese. Darbdavys galėtų susidurti su teisinių pasekmių rizika nuo Lygių galimybių kontrolieriaus tyrimo iki skundų pagal BDAR ar net civilinės atsakomybės bylų. Todėl DI taikymas įdarbinimo procese suponuoja iš esmės naujas darbdavio atsakomybes bei pareigą užtikrinti, kad technologijos nepažeistų žmogaus teisių.

Minėtos airės, AI aktas taip pat patvirtina, kad tikralaikio biometrinio atpažinimo viešose vietose taikymas leidžiamas tik išimtiniais atvejais, susijusiais su labai sunkiais nusikaltimais, ir tik gavus išankstinį teismo leidimą. Tai svarbu ir darbo aplinkoje, ypač kai įmonės veikia viešose ar pusiau viešose erdvėse pvz., prekybos centruose, transporto sektoriuje ar klientų aptarnavimo vietose. Praktikoje tai reiškia, kad automatinis veidų atpažinimas, emocijų analizė ar elgesio sekimas realiu laiku darbo vietose gali būti laikomi nepriimtinos rizikos DI sistemomis, kurių naudojimas nuo 2025 m. vasario 2 d. yra griežtai draudžiamas. Tokie įrankiai turi būti nutraukti arba pritaikyti prie DI akto reikalavimų. Tokiu būdu Europos teisės sistema riboja darbdavio technologinę galią ir užtikrina, kad technologinis progresas neperžengtų darbuotojo autonomijos, teisės į privatų gyvenimą ribų.

Dirbtinio intelekto technologijų diegimas darbo santykiuose privalo būti vertinamas per asmens pagrindinių teisių apsaugos prizmę. Bendrasis duomenų apsaugos reglamentas (BDAR) aiškiai nustato, kad asmuo turi teisę nebūti sprendimo, kuris pagrįstas vien tik automatizuotu duomenų tvarkymu, objektu, jeigu toks sprendimas sukelia teisinius padarinius arba panašiai reikšmingai paveikia jo padėtį (BDAR, 22 str. 1 d.). Tai leidžia daryti išvadą, kad sprendimai, priimami remiantis vien DI algoritmais pavyzdžiui, dėl priėmimo į darbą, vertinimo ar atleidimo – negali būti teisėti, jei juose visiškai nėra žmogaus įsikišimo, išskyrus tris aiškiai reglamentuotus atvejus: kai sprendimas būtinas sutarčiai vykdyti, kai leidžia įstatymas, arba kai darbuotojas duoda aiškų ir informuotą sutikimą (BDAR, 22 str. 2 d.).

Darbo teisės kontekste šis principas yra ypač reikšmingas, nes DI sprendimai dažnai taikomi būtent tokiose situacijose, kurios daro reikšmingą įtaką darbuotojo teisei padėčiai, pavyzdžiui, vertinant našumą ar priimant sprendimus dėl atleidimo.. Tokiais atvejais vien technologinis sistemos efektyvumas ar tikslumas negali pateisinti asmens teisių apribojimo, nes pagal BDAR 22 straipsnį asmuo turi teisę į žmogaus įsikišimą, nuomonės išreiškimą ir sprendimo paaiškinimą. Tai reiškia, kad technologinis progresas privalo būti derinamas su teisiniais principais: teisėtumu, skaidrumu ir paaiškinamumu, kurie įtvirtinti BDAR ir DI reguliavimo dokumentuose bei pabrėžiami Europos Komisijos etikos gairėse (2019).Galime

teigti, kad DI iš esmės kvestionuoja klasikinę darbo teisės paradigmą, paremtą subordinacijos, kontrolės ir priklausomybės principais. Kai DI algoritmai įgyja galimybę spręsti dėl darbo užduočių paskirstymo, pamainų planavimo ar net darbuotojų vertinimo, tradicinės darbdavio funkcijos pradeda skaidytis tarp juridinio darbdavio ir algoritminės sistemos. Tokia tendencija kuria prielaidas atsirasti vadinamajam „algoritminiam valdymui“ (angl. *algorithmic management*) – reiškiniui, kuris, pasak De Stefano, iš esmės transformuoja darbo santykių prigimtį ir reikalauja naujo teisinio reguliavimo modelio (De Stefano ir Taes, 2022).

Kaip pastebi Tamošaitis ir Gaubienė (2025), dirbtinis intelektas jau dabar diegiamas teisėsaugos institucijose – nuo automatizuotų stebėjimo sistemų iki duomenų analizės kriminalistiniuose tyrimuose. „Nors tokios sistemos gali padidinti operatyvumą, kyla rimtų klausimų dėl neteisėto asmens duomenų naudojimo, klaidingo kaltinimų paskelbimo ar net diskriminacijos rizikos“ (Tamošaitis ir Gaubienė, 2025). Jie išskiria vieną iš esminių problemų – algoritminį šališkumą, kai DI modeliai mokomi su neobjektyviais arba istoriškai diskriminaciniais duomenimis ir tokiu būdu pakartoja socialinius netolygumus.

Algoritminė diskriminacija – vienas reikšmingiausių iššūkių dirbtinio intelekto sistemų taikyme darbo ar viešojo administravimo srityje. Kaip pažymi Wachter, Mittelstadt ir Floridi (2017), algoritmų pagrindu priimti sprendimai dažnai remiasi istoriniais duomenimis, kurie gali būti iš prigimties šališki, taip reprodukuodami socialinius neteisingumus. Tai ypač pavojinga, kai DI taikomas sprendimams, darančiais poveikį asmens teisėms – atrankai į darbą, veiklos vertinimui ar socialinėms paslaugoms. Autoriai pabrėžia, kad teisė į sprendimo paaiškinimą dažnai lieka formali, ypač kai sprendimų algoritminė logika nėra suprantama net darbdaviui, jei nėra aiškiai apibrėžta, kokie duomenys naudoti, kaip modelis veikia ir kas už jį atsako.

Panašiai ir Ineta Breskienė (2024) nurodo, kad algoritminio valdymo sistemos darbo teisėje dažnai veikia kaip „juodosios dėžės“, kurių veikimo logika nepermatoma nei darbuotojui, nei darbdaviui. Tai sukuria teisinį vakuumą, kuriame darbuotojas lieka be galimybės išsiaiškinti, kodėl jis buvo neigiamai įvertintas ar atleistas, o darbdavys – be realių priemonių įvertinti sistemos teisėtumą.

Todėl būtina inicijuoti teisės aktų pokyčius, aiškiai reglamentuojančius DI taikymą darbuotojų vertinimo procese. Lietuvos Respublikos darbo kodekse turėtų būti įtvirtinta, kad automatizuotas sprendimų priėmimas darbo aplinkoje yra leidžiamas tik laikantis pagrindinių principų: darbuotojo informavimo, sprendimo paaiškinimo, žmogaus įsikišimo ir apskundimo galimybės. Šios nuostatos galėtų būti įtrauktos kaip atskiras DK straipsnis, reglamentuojantis dirbtinio intelekto taikymą darbo santykiuose. Be to, darbdaviai turėtų nustatyti vidines procedūras, reglamentuojančias DI sistemų veikimą, jų logikos suprantamumą bei sprendimų

skaidrumą. Taip būtų užtikrinamas ne tik technologinis efektyvumas, bet ir pagrindinių darbuotojo teisių įgyvendinimas.

1.3. Lietuvos Respublikos darbo kodeksas ir jo santykis su DI reguliavimu

Lietuvos Respublikos darbo kodeksas yra pagrindinis darbo santykius reglamentuojantis teisės aktas, tačiau jame iki šiol nėra tiesiogiai aptariamos dirbtinio intelekto sistemos ar jų taikymas darbo teisiniuose santykiuose. Vis dėlto, esami principai, tokie kaip proporcingumo, nediskriminavimo bei skaidraus darbo organizavimo, leidžia interpretuoti DI taikymo ribas bei keliamus reikalavimus darbo aplinkoje.

Lietuvos Respublikos darbo kodekso 206 straipsnyje nustatyta darbdavio pareiga konsultuotis su darbuotojų atstovais atsiranda tik tais atvejais, kai priimami ar keičiami specifiniai vietiniai norminiai teisės aktai, tiesiogiai susiję su šiame straipsnyje išvardintomis sritimis. Dirbtinio intelekto diegimas darbo vietoje savaime nesukuria tokios pareigos, tačiau ji gali kilti, jeigu DI taikymo aspektai yra įforminami dokumentais, reglamentuojančiais, pavyzdžiui, darbuotojų stebėseną (206 str. 5 p.), privatumo pažeidimo rizikas (6 p.) ar darbuotojų asmens duomenų apsaugą (7 p.). Tokiu atveju darbdavys turi laikytis konsultavimosi procedūrų, nes sprendžiami klausimai dėl įmonės vidinių norminių aktų priėmimo ar keitimo, susijusių su DI reguliavimu darbo teisiniuose santykiuose, patenka į 206 straipsnio taikymo apimtį. Vis dėlto, jei DI sistemos įdiegiamos darbo aplinkoje nepriimant jokių naujų vidaus teisės aktų ar taisyklių, konsultacijų pareiga pagal DK 206 str. formaliai netaikoma.

Darbo kodekso 203 straipsnyje įtvirtinta darbuotojų ir jų atstovų teisė į informavimą ir konsultavimą darbdavio sprendimų, susijusių su darbuotojų darbo, socialinių, ekonominių teisių bei interesų įgyvendinimu ir gynimu, klausimais. Jei DI diegimas daro poveikį minėtoms darbuotojo teisėms, galima teigti, kad teisė į informavimą ir konsultavimą privalo būti įgyvendinta.

Pagal Bendrojo duomenų apsaugos reglamento (BDAR) 13–15 straipsnius, darbdavys, tvarkantis darbuotojo asmens duomenis, privalo pateikti informaciją apie duomenų tvarkymo tikslus, pagrindą, trukmę, duomenų šaltinius, gavėjus ir, svarbiausia – apie tai, ar taikomas automatizuotas sprendimų priėmimas, įskaitant profiliavimą. Pavyzdžiui, jeigu darbdavys naudoja DI sistemą, kuri analizuoja darbuotojų našumo rodiklius arba vertina darbuotojų elgseną, tačiau apie tai darbuotojai nėra aiškiai informuoti – pažeidžiami BDAR informavimo reikalavimai. Tokia informacija gali būti pateikiama abstrakčiai vidaus tvarkose ar paslėpta techniniuose aprašymuose, todėl darbuotojai gali ir nesuprasti, kaip jų duomenys naudojami ar vertinami pasitelkiant DI. Tokios situacijos pažeistų darbuotojų teises.

Dirbtinio intelekto akte (Reglamentas (ES) 2024/1689, 29 str.) papildomai įtvirtinta darbdavio pareiga informuoti darbuotojus apie DI sistemų veikimą, jų paskirtį, galimą poveikį

sprendimams bei užtikrinti, kad ši informacija būtų prieinama, aiški ir suprantama. Tačiau jei darbdavys apsiriboja tik vidaus taisyklių paskelbimu, kurios nėra išsamios ar suprantamos, toks „informavimas“ tampa formaliu, neatitinkančiu nei BDAR, nei DI akto tikslų. Todėl net jei konsultacijų pareiga pagal DK 206 straipsnį neatsiranda, darbdavys negali išvengti pareigos užtikrinti skaidrumą, informavimą ir darbuotojų teisių apsaugą – ne tik formaliai, bet ir turinio prasme, tam taip pat turėtų būti taikomas DK 203 straipsnis.

DK 26 str. užtikrina lygias galimybes ir draudžia diskriminaciją, o DI sistemų sprendimai, pagrįsti netiesioginiu šališkumu ar istoriniais duomenų šaltiniais, gali tapti diskriminacijos šaltiniu. Tokiu atveju, darbdaviui gali kilti pareiga pagrįsti, jog DI formuojami sprendimai buvo neutralūs ir proporcingi siekiamam tikslui, remiantis ir BDAR nuostatomis bei DI akto reikalavimais (Reglamentas (ES) 2024/1689).

Darbo sutarties pasibaigimo ir nutraukimo pagrindai bei tvarka reglamentuojami Lietuvos Respublikos darbo kodekso 54–64 straipsniuose. Tačiau DK nenumato specifinių reikalavimų dėl sprendimų priėmimo, kai šie grindžiami dirbtinio intelekto sistemomis. Jei darbo sutarties nutraukimo sprendimas priimamas remiantis algoritminiu darbuotojo vertinimu ar automatizuotais našumo rodikliais, kyla klausimas dėl sprendimo pagrįstumo, žmogaus įsikišimo ir kaip jį tikslingai ginčyti, jei neturima algoritminio sprendimo logikos paaiškinimo. Remiantis Bendrojo duomenų apsaugos reglamento 22 straipsniu, asmuo turi teisę nebūti vien automatizuoto sprendimo, sukeliančio reikšmingas pasekmes, subjektu. Be to, pagal Europos Parlamento ir Tarybos reglamentą (ES) 2024/1689 dėl dirbtinio intelekto, tokie sprendimai priskiriami prie „didelės rizikos“ sistemų, kurios privalo atitikti skaidrumo, paaiškinamumo, žmogaus įsikišimo ir atsakomybės principus. Lietuvos darbo kodeksas šiuo metu tiesiogiai nereglamentuoja būtent DI sprendimų reguliavimo darbo sutarties nutraukimo atvejais, nebent galėtų būti taikomas DK 203 straipsnis, bet siekiant išvengti netikslumų, svarstyтина galimybė DK papildyti konkrečiomis nuostatomis dėl sprendimų priėmimo taikant DI arba procedūrinių reikalavimų konkrečiai ir darbo sutarčių nutraukimo atvejais. Nuostatos turėtų įtvirtinti darbuotojo teisę būti informuotam apie DI sprendimą, reikalauti žmogaus peržiūros, tik taip užtikrinamas teisinis saugumas technologinės transformacijos laikotarpiu.

Darbdavio pareiga užtikrinti darbuotojų teisę į privatų gyvenimą bei asmens duomenų apsaugą kyla iš Lietuvos Respublikos darbo kodekso 27 straipsnio, taip pat iš Bendrojo duomenų apsaugos reglamento ir Lietuvos Respublikos asmens duomenų teisinės apsaugos įstatymo.

Mokslinėje literatūroje vis dažniau akcentuojama būtinybė peržiūrėti DK nuostatas, atsižvelgiant į Europos Sąjungos teisės raidą, ypač į Dirbtinio intelekto aktą. Dirbtinio intelekto diegimas darbo santykiuose daro tiesioginį poveikį darbuotojų teisėms – nuo vertinimo iki atleidimo. Kadangi šios situacijos nėra detalios reglamentuotos BDAR ar ES DI akte, reikalinga

papildyti Lietuvos Respublikos darbo kodeksą specialiomis nuostatomis. Nacionalinis reglamentavimas turi užtikrinti ne tik bendrą atitiktį ES teisės aktams, bet ir sudaryti sąlygas įgyvendinti darbo teisės principus – proporcingumą ir skaidrumą – konkrečiose praktinėse situacijose, kurias sukelia algoritminis valdymas. Tai apima pareigą informuoti darbuotojus apie DI veikimą, sprendimų logiką, numatyti žmogaus įsikišimą ir detalizuoti apskundimo galimybę. Kaip pažymi Ineta Breskienė (2024), algoritminio valdymo praktika kelia didžiausią grėsmę ne vien dėl sprendimų turinio, bet dėl jų priėmimo proceso neskaidrumo – darbuotojas dažnai neturi galimybės suprasti, kokie duomenys buvo naudoti, kaip sprendimas buvo suformuotas ar jis turi teisę jį ginčyti. Autorė pabrėžia, kad tokios sistemos pažeidžia darbuotojo autonomiją, privatumo apsaugą bei teisę į sprendimo pagrįstumą, ypač kai darbdavys remiasi komercinėmis, uždaros struktūros DI platformomis, veikiančiomis kaip vadinamosios „juodosios dėžės“. Ši problema tampa itin aktuali, kai sprendimai dėl atrankos, vertinimo ar net atleidimo priimami vien remiantis profiliavimu, elgesio analize ir jeigu jie priimami be tiesioginio žmogaus dalyvavimo.

Tuo tarpu Donatas Murauskas (2020) siūlo konstruktyvų sprendimą: teisiškai klasifikuoti algoritmus pagal jų taikymo stadijas sprendimų priėmimo procese. Jis išskiria tris stadijas – faktinių aplinkybių analizę – algoritmai naudojami duomenims apdoroti, vertinimams atlikti, bet ne priimti galutiniams sprendimams; taikytinos teisės parinkimą – algoritmai padeda surasti aktualius norminius šaltinius ar taikyti teisminę praktiką; Galutinį sprendimo priėmimą – algoritmai gali rekomenduoti ar prognozuoti rezultatą, tačiau galutinis sprendimas turi būti priimamas žmogaus. Tokia klasifikacija leistų diferencijuoti reguliavimo intensyvumą: rekomendaciniai algoritmai (pvz., pagalbinės analitinės priemonės) galėtų būti reguliuojami švelniau, tuo tarpu autonominiai sprendimų priėmimo modeliai (pvz., automatinis kandidatų eliminavimas iš atrankos) turėtų būti griežčiau vertinami tiek BDAR, tiek pagal AI aktą. Asistuojančioji (tarpinė) stadija turėtų būti vertinama pagal tai, kiek į procesą įtraukiamas žmogaus sprendimas, paaiškinamumas ir apskundimo galimybės. Tai reikštų, kad teisinio reglamentavimo intensyvumas būtų proporcingas algoritminės galios poveikiui žmogaus teisėms.

Abi šios pozicijos išryškina esminį darbo teisės iššūkį DI eroje – kaip užtikrinti, kad technologinis efektyvumas nenuvertintų procesinio teisingumo, darbuotojų dalyvavimo, konsultavimo, informavimo ir žmogaus įsikišimo kiekvienu atveju ir teisės į sprendimų paaiškinimą. Tai ypač svarbu kontekste, kai algoritmai padeda priimti sprendimus dėl darbo užduočių paskirstymo, premijų ar net atleidimo. Šiuo požiūriu, Murausko pasiūlyta klasifikacija galėtų būti integruota į AI akto taikymą praktikoje, suteikiant galimybę atskirti tikslines ir rizikingas sritis, kuriose būtina teisės subjekto apsauga, nuo tų, kur DI veikia tik kaip pagalbinė priemonė.

Jeigu darbdavys remiasi automatizuotomis sistemomis, vertindamas darbuotojo našumą ar paskirstydamas premijas, būtina užtikrinti, kad darbuotojai būtų informuoti apie sprendimų logiką (Reglamentas (ES) 2016/679).

Europos socialinės chartijos 21 straipsnis įtvirtina darbuotojų ar jų atstovų teisę būti informuotiems ir konsultuojamiems apie sprendimus, galinčius turėti reikšmingą poveikį jų darbo sąlygoms ar užimtumo padėčiai įmonėje, ypač ekonominių pertvarkymų ar darbo organizavimo pokyčių kontekste. Nors ši nuostata tiesiogiai nemini dirbtinio intelekto, jos taikymo principas gali būti aktualus ir DI technologijų diegimo atvejais. Jei DI sprendimai daro poveikį darbo sąlygoms, profesinėms funkcijoms ar jų apimčiai, darbuotojų atstovai turėtų turėti galimybę dalyvauti sprendimų priėmimo procese (Europos socialinė chartija, 1996). Šis konstitucinio lygmens dokumentas įpareigoja valstybes nares sudaryti sąlygas realiam informavimo ir konsultavimo įgyvendinimui, todėl verta būtų apsvarstyti teisės aktų pritaikymą konkrečiai DI situacijoms dėl informavimo ir konsultavimo procedūrų.

Lietuvos Respublikos darbo kodeksas aiškiai apibrėžia darbdavio teisę nutraukti darbo sutartį dėl perteklinių funkcijų (DK 57 straipsnis), tačiau dirbtinio intelekto diegimas gali sukurti naują situaciją: kai funkcija tampa perteklinė ne dėl objektyvaus darbo poreikio sumažėjimo, o dėl technologinės galimybės ją automatizuoti. Tokiu atveju kyla klausimas, ar atleidimas, grindžiamas vien automatizavimo galimybe, iš tikrųjų atitinka DK 57 straipsnio prasmę. Teisiškai žiūrint, jei darbo funkcija yra panaikinama ir nebeatliekama žmogaus, tai gali būti traktuojama kaip pagrindas atleisti darbuotoją be kaltės dėl perteklinės darbo funkcijos. Žinoma, kilus ginčui turi būti atsižvelgiama į realią situaciją ir susiklosčiusias aplinkybes, vertinant, ar funkcija tikrai tapo pertekline. Kaip pažymi De Stefano (2020), algoritminio valdymo plėtra keičia pačią darbo netekimo logiką – darbuotojas atleidžiamas ne dėl to, kad darbas išnyksta, bet todėl, kad jį gali atlikti automatizuota sistema. Tai kelia klausimų dėl socialinio teisingumo, technologijų kontrolės ir darbuotojų apsaugos. Toks požiūris gali skatinti piktnaudžiavimą automatizavimu kaip atleidimo priemone, todėl būtina svarstyti galimybę darbo teisėje įtvirtinti papildomas garantijas – pavyzdžiui, reikalavimą įrodyti, kad automatizavimas yra būtinas, arba pareigą pasiūlyti perkvalifikavimą. Šiuo požiūriu DK 57 straipsnio taikymas turėtų būti aiškinamas sistemiškai, kartu su socialinės atsakomybės ir technologijų etiško diegimo principais.

Sprendžiant ginčus dėl darbo santykių pokyčių, inicijuotų dirbtinio intelekto taikymo, būtina remtis ne tik Darbo kodekso nuostatomis, bet ir esminiais darbo teisės principais, įtvirtintais DK 24 straipsnyje – proporcingumo, socialinio dialogo, sąžiningumo ir nediskriminavimo principais. Šie principai įgyja naują reikšmę, kai sprendimus dėl darbo vietos, vertinimo ar atleidimo nebe vien vadovas, o tai daroma su algoritmo pateiktu vertinimu,

nuomone. Tokiu atveju žmogaus orumo apsauga, įtvirtinta ir Lietuvos Respublikos Konstitucijos 21 straipsnyje, suponuoja, kad net technologinis sprendimas turi būti priimamas etiškai, skaidriai ir proporcingai, o darbuotojas turi būti traktuojamas kaip savarankiška sprendimo proceso šalis, turinti teisę suprasti, reaguoti ir dalyvauti.

Europos Komisija siūlydama Platforminio darbo direktyvos projektą (Platforminio darbo direktyvos pasiūlymas, COM(2021) 762 final) aiškiai nurodė būtinybę įtvirtinti darbuotojų teisę žinoti apie algoritminio valdymo priemones, jų logiką ir poveikį darbo rezultatų vertinimui, atlygiui ar grafikui. Direktyvos 6 straipsnis aiškiai įpareigoja darbdavį užtikrinti informavimo, konsultavimo bei žmogaus įsikišimo garantijas. Tokia kryptis rodo tendenciją, kad DI tampa ne tik technologine naujove, bet ir darbo santykių reguliavimo objektu.

Šios plėtros kryptys leidžia formuluoti esminę darbo teisės aktualiją: DI nėra paprastas techninis pokytis – tai struktūrinis darbo teisės iššūkis, kuris reikalauja adaptacijos tiek nacionaliniu, tiek Europos lygiu. Kaip pažymi Tamošaitis ir Gaubienė (2025), DI tampa precedento neturinčia priemone, kuri gali padidinti produktyvumą, tačiau kartu kelia klausimų dėl atsakomybės ir darbuotojų teisių apsaugos. Šiuo požiūriu darbo teisė, kaip socialinio teisingumo ir galios pusiausvyros palaikymo sistema, turi būti rekonstruojama taip, kad jos apsauga apimtų ir algoritminio valdymo sritis – net jei DI sprendimai yra techninio pobūdžio, jie daro realų poveikį žmogaus teisei padėčiai (De Stefano, 2020).

Darbo teisė iš esmės buvo kuriama tam, kad apsaugotų žmogų, atliekantį darbą savo pastangomis, sąveikaujantį su darbdaviu tiesioginiuose socialiniuose santykiuose. Tokį darbą visada lydėjo atsakomybė, sprendimų logika ir asmeninis kontaktas tarp darbo teisės subjektų. Tačiau šiuolaikinėje darbo aplinkoje šią santykio ašį vis labiau keičia dirbtinio intelekto pagrindu veikiančios sistemos, kurios koreguoja darbo organizavimo, kontrolės ar vertinimo funkcijas. Tokiu kontekstu tampa būtina atsigręžti į esminius darbo teisės principus. Kaip nurodo Bitinas, Tartilas ir Litvaitienė (2011), darbo teisė grindžiama orumo, lygiateisiškumo bei apsaugos nuo socialinės rizikos principais, kurie išlieka aktualūs net ir kintant darbo turiniui (p. 291). Nors autoriai šių principų netaiko dirbtinio intelekto situacijoms, šios vertybinės nuostatos įgyja naują aktualumą algoritminio darbo kontekste, kur tradicinės žmogaus ir darbdavio sąveikos formos iš esmės keičiasi.

Kai sprendimus dėl darbuotojo – jo darbo rezultatų, atlygio ar net atleidimo pradeda priimti nebe tiesiogiai pats vadovas, o remiamasi algoritminės sistemos vertinimu ar pagalba, keičiasi pats pavaldumo santykis ir darbo teisinės atsakomybės struktūra. Darbuotojas dažnai neturi galimybės suprasti, kas faktiškai atsakingas už sprendimą, kokie kriterijai taikyti, kokia sprendimo logika. Tokia situacija apibūdinama kaip „algoritminis valdymas“ – kai sprendimų priėmimo galia *de facto* perleidžiama programinei įrangai, kuri veikia nepriklausomai nuo

tiesioginės žmogaus kontrolės (De Stefano ir Taes, 2022). Susidaro teisinė dilema: nors atsakomybė už sprendimus formaliai tenka darbdaviui, praktikoje sprendimų priėmimo procesas pasiskirsto tarp algoritmų, platformų tiekėjų ir technologijų kūrėjų. Toks darbo organizavimas kelia naujų klausimų dėl to, kas iš tikrųjų atsakingas, kaip sprendimai priimami ir ar darbuotojas turi galimybę suprasti bei dalyvauti šiame procese.

Jei sprendimas dėl darbuotojo atleidimo ar vertinimo priimamas automatizuotai, darbuotojas turi teisę žinoti, kaip veikia sistema, kas už tokį sprendimą atsakingas. Kaip pažymi Breskienė (2024), tokie sprendimai kelia riziką darbuotojo autonomijai ir privatumo apsaugai, nes „darbuotojas paliekamas vienas prieš sudėtingas, komerciškai uždaras technologines sistemas“ (p. 114). Todėl būtina iš naujo įvertinti, kaip darbo teisė gali užtikrinti žmogaus apsaugą ne tik nuo darbdavio veiksmų, bet ir nuo algoritmizuoto valdymo, kuris formaliai atrodo objektyvus, bet faktiškai gali būti diskriminuojantis ar nepaaiškinamas. Todėl svarbus darbdavys žmogus, kaip galutinis sprendimo priėmėjas.

Pagal Dirbtinio intelekto akto (Reglamentas (ES) 2024/1689) I priedo 4(a) punktą, DI sistemos, taikomos darbo santykiuose – įskaitant sprendimus dėl įdarbinimo, vertinimo, užduočių paskirstymo ar atleidimo – laikomos didelės rizikos sistemomis. Todėl joms taikomi specialūs reikalavimai: sprendimų paaiškinamumas, žmogaus įsikišimo užtikrinimas, dokumentavimas ir rizikos vertinimas (29 str.). Tokia pozicija grindžiama ne tik ES teisės aktų logika, bet ir šiuolaikinėmis nacionalinėmis mokslinėmis įžvalgomis. Grumulaitis (2023) dar labiau pabrėžia, kad darbuotojas negali būti laikomas atsakingu už sprendimus, priimtus algoritminės sistemos, jei jam nebuvo suteikta informacija apie jos veikimo principą ar jis neturėjo realios galimybės į tai įsikišti.

Automatizuoti sprendimai dažnai grindžiami asmens veiklos, elgsenos, emocinės būklės ir net fizinio judėjimo duomenimis. Tokie sprendimai patenka į Bendrojo duomenų apsaugos reglamento (BDAR) 22 straipsnio taikymo sritį, kuris suteikia duomenų subjektams teisę nebūti sprendimo, grindžiamo tik automatizuotu duomenų tvarkymu, subjektu, jei toks sprendimas sukelia teisinės pasekmės arba panašiai reikšmingai juos veikia. Tačiau, kaip pažymi Wachter, Mittelstadt ir Floridi (2017), ši teisė dažnai yra ribota praktikoje, nes BDAR reikalavimai dėl informacijos suteikimo apie automatizuotą sprendimų priėmimą apsiriboja bendro pobūdžio informacija apie taikomą logiką, o ne konkrečiais paaiškinimais apie individualius sprendimus. Dėl to darbuotojai gali nežinoti, kad jų veikla yra vertinama automatizuotai, arba nesuprasti, kaip tokie sprendimai yra priimami, kas apsunkina galimybę ginčyti tokius sprendimus. Todėl svarbus darbuotojų informavimas apie DI sistemų naudojimą.

Dirbtinio intelekto sistemų diegimas gali turėti reikšmingą įtaką darbo turiniui, ypač tais atvejais, kai DI tampa atsakingas už darbo paskirstymą, veiklos vertinimo kriterijų taikymą ar

sprendimų dėl užduočių ir atlygio priėmimą. Praktikoje gali susidaryti situacijų, kai darbuotojo pareigos formaliai lieka nepakitusios, tačiau faktinė jų įgyvendinimo struktūra keičiasi iš esmės. Pavyzdžiui, darbuotojas nebesprendžia, ką ir kada atlikti, nes tai nustato DI sistema pagal algoritminį modelį, o jo veiklos rezultatai vertinami pagal automatizuotus rodiklius. Tokie pokyčiai kelia klausimą, ar tai tebėra darbdavio teisė vienašališkai keisti darbo organizavimą (kaip leidžiama pagal DK 35 ir 36 straipsnius), ar jau yra darbo funkcijų, kaip būtinosios darbo sutarties sąlygos, pasikeitimas, kuris pagal DK 34 ir 45 straipsnius reikalauja rašytinio darbuotojo sutikimo.

Pavyzdžiui, jeigu prekybos įmonėje dirbantis vadybininkas anksčiau pats planuodavo užduotis, bet po DI diegimo šias užduotis automatiškai paskirsto sistema, o darbuotojas tampa tik algoritminės logikos vykdytoju ar stebėtoju, tai lemia ne tik darbo pobūdžio, bet ir funkcijų perorientavimą, nuo savarankiško sprendimų priėmimo prie kontrolinio, prižiūrinčio vaidmens. Kadangi pagal DK 33 ir 34 straipsnius darbo funkcijos yra aiškiai apibrėžta darbo sutarties dalis, tokio tipo pokytis turėtų būti vertinamas kaip darbo sąlygų keitimas, o ne tik organizacinė pertvarka. Vis dėlto nacionalinėje teisėje šiuo metu nėra aiškiai reglamentuota, kada technologinio sprendimo įvedimas laikytinas sutartinių sąlygų keitimu, todėl egzistuoja teisinio neapibrėžtumo rizika. Siekiant aiškumo, būtų tikslinga svarstyti Darbo kodekso papildymą, nustatant kriterijus, taip pat, kur yra ta riba, pagal kurią, funkcijų automatizavimas ar darbuotojo vaidmens pasikeitimas į DI priežiūrą laikytinas darbo funkcijų keitimu, reikalaujančiu darbo sutarties pakeitimo. Toks norminis tikslinimas užtikrintų darbuotojų teisių apsaugą besikeičiančioje technologinėje darbo aplinkoje ir pašalintų interpretacinę pilkąją zoną.

Kaip pažymi Grumulaitis (2023), „funkciniai darbo pokyčiai, atsirandantys dėl algoritminio valdymo, ne visada yra aiškiai matomi darbo sutartyse, tačiau jų pasekmės darbuotojui yra esminės – keičiasi jo kontrolės laipsnis, veiklos vertinimo kriterijai, o kartais ir atsakomybės mastas“ (p. 5–6). Tai reiškia, kad net jei formaliai sutarties tekstas lieka nepakitęs, faktiniai darbo santykiai gali būti transformuoti tiek, kad prilygtų esminėms sąlygoms. Tokiais atvejais, nesant darbuotojo sutikimo, būtų pagrįsta abejoti, ar nepažeidžiamas 45 straipsnio reikalavimas dėl sąlygų keitimo tik su darbuotojo sutikimu. Šiame kontekste darbo teisės normų taikymas turi būti nukreiptas ne tik į formalių požymių analizę, bet ir į turinio vertinimą – ar faktiniai pokyčiai atitinka darbo teisiniams santykiui keliamus skaidrumo, lojalumo ir informavimo principus.

Skirtingai nei Europos Sąjungos dirbtinio intelekto aktas (Reglamentas (ES) 2024/1689), kuris įtvirtina naudotojo pareigą dokumentuoti sprendimų logiką, užtikrinti skaidrumą ir žmogaus įsikišimą į sprendimų priėmimą (29 str.). Lietuvos Respublikos darbo kodekse šiuo metu nėra aiškių taisyklių, kaip turi būti pagrįstas ar peržiūrimas sprendimas, priimtas naudojant

DI sistemą. Nors Darbo kodeksas apima bendrąsias darbdavio pareigas užtikrinti darbuotojų teisių apsaugą, informuotumą, jis nenumato specifinių mechanizmų, kaip, kada darbuotojas turėtų būti informuojamas apie algoritmų priimtus sprendimus, kaip šie sprendimai turėtų būti motyvuojami ar koku būdu darbuotojas galėtų reikalauti jų peržiūros. Tai gali sukurti teisinę spragą, ypač atleidimo, atrankos ir našumo vertinimo atvejais, kai sprendimo automatizavimas gali iš esmės paveikti darbuotojo teises. Tuo tarpu ES teisėje šie klausimai yra aiškiau reglamentuoti. Pagal Bendrojo duomenų apsaugos reglamento (BDAR) 22 straipsnį, jei sprendimas yra priimamas vien automatizuotai ir turi reikšmingą poveikį darbuotojui (pvz., atleidimas, vertinimas), darbuotojas turi teisę: i) nebūti vien automatizuoto sprendimo objektu, ii) gauti informaciją apie tokio sprendimo logiką, iii) reikalauti žmogaus įsikišimo bei ginčyti sprendimą. Šias nuostatas papildomai konkretina Dirbtinio intelekto akto 29 straipsnis, kuris įpareigoja darbdavius informuoti darbuotojus apie DI sistemų paskirtį, veikimo principus ir poveikį, taip pat užtikrinti, kad sprendimai galėtų būti suprantami ir paaiškinami.

Darbdavio pareigos taikyti šiuos reikalavimus nėra aiškiai konkretizuotos nacionalinės teisės aktuose, kurie reglamentuoja darbdavių ir darbuotojų darbo teisinius santykius. Nacionalinis reglamentavimas darbo kodekse galėtų ne tik sustiprinti darbuotojų teisių apsaugą, bet ir užtikrinti teisinį tikrumą bei aiškesnę atsakomybės sistemą, palengvinant ginčų nagrinėjimą nacionalinėje teisminėje praktikoje. Priešingu atveju, priklausomybė nuo ES reglamentų be nacionalinio įgyvendinimo gali sukelti praktinio taikymo iššūkių. Be aiškaus reglamentavimo Darbo kodekse, BDAR ir AI akto nuostatos gali būti taikomos fragmentiškai, skirtingų darbdavių ar sektorių praktikoje nevienodai. Nacionalinis reglamentavimas Darbo kodekse galėtų sustiprinti darbuotojų teisių apsaugą, įnešti aiškumo tiek patiems darbdaviams ir darbuotojams dėl konkrečių normų DI atveju taikymo, užtikrinti vienodą normų taikymą praktikoje ir palengvinti ginčų sprendimą. Tai taip pat aiškiau apibrėžtų darbdavių atsakomybę ir sumažintų riziką, kad informavimo ir sprendimų pagrįstumo pareigos bus vykdomos tik formaliai arba visai ignoruojamos.

Paradoksalu, tačiau moderniausios algoritminės technologijos dažnai veikia remiantis teisinėmis taisyklėmis, kurios buvo sukurtos kitame istoriniame ir technologiniame kontekste – kai žmogaus sprendimas, o ne automatinis duomenų apdorojimas, buvo laikomas atsakomybės epicentru. Toks neatitikimas sukuria vadinamąjį reguliacinį disonansą: teisinė atsakomybė lieka žmogui, bet sprendimo logika priklauso algoritminėms sistemoms, veikiančioms be aiškos prieigos ar paaiškinimo. Nors AI aktas (Reglamentas (ES) 2024/1689) reikalauja skaidrumo, dokumentavimo ir žmogaus įsikišimo didelės rizikos DI srityse (29 str.), faktinė technologinė praktika gali to ir neatspindėti. Ypač tada, kai naudojamos trečiųjų šalių komercinės platformos, kurių veikimo principai yra intelektinė nuosavybė ir gali būti neatskleidžiami.

Nors dirbtinio intelekto integravimas į darbo santykius pirmiausia siejamas su teisiniais įsipareigojimais, vis aktualesniu tampa klausimas, kaip šiuos procesus įgyvendinti etiškai ir socialiai atsakingai. Europos Komisijos parengtos Patikimo DI gairės (2019) pabrėžia, kad skaidrumas, žmogaus kontrolė, paaiškinamumas ir atskaitomybė turi būti laikomi ne techninėmis, o vertybinėmis sąlygomis bet kokiam DI taikymui. Tai leidžia daryti išvadą, kad sprendimų autonomizavimas neturi eliminuoti žmogaus įsitraukimo bei atsakomybės. Kaip pažymi De Stefano ir Aloisi (2022), technologinių transformacijų sėkmė priklauso nuo to, kiek šie pokyčiai yra derinami su kolektyviniu socialiniu dialogu, kuris užtikrina darbuotojų dalyvavimą bei įtaką formuojant algoritmines praktikas. Todėl DI negali būti traktuojamas tik kaip modernizacijos įrankis – jis tampa ir socialinio kontrakto tarp darbdavio ir darbuotojo išplėtimo objektu.

2. DIRBTINIO INTELEKTO POVEIKIS DARBUOTOJŲ TEISĖMS

2.1. Biometrinių duomenų naudojimas ir privatumo teisės apsauga

Dirbtinio intelekto technologijų diegimas darbo aplinkoje ypač išryškina su biometrinių duomenų apsauga ir darbuotojo privatumu susijusias teises ribas. Lietuvos Respublikos darbo kodekso 27 straipsnyje įtvirtinta darbuotojo teisė į privatą gyvenimą darbo vietoje, įskaitant apsaugą nuo perteklinės darbdavio kontrolės ar kišimosi į asmeninę erdvę. Ši teisė glaudžiai siejasi su darbuotojų asmens duomenų apsaugos principais, kuriuos darbdavys privalo užtikrinti taikydamas informacines technologijas, ypač tokias, kurios gali apdoroti asmens tapatybės duomenis ar stebėti darbuotojų veiklą. Ši nuostata kartu su BDAR 5 ir 6 straipsnių reikalavimais dėl duomenų teisėtumo ir proporcingumo riboja DI sistemų naudojimo galimybes darbo aplinkoje. Tuo tarpu DK 206 straipsnyje nustatoma darbdavio pareiga konsultuotis su darbuotojų atstovais tik tais atvejais, kai rengiamas ar keičiamas vietinis norminis aktas, susijęs, be kita ko, su darbuotojų stebėseną ar asmens duomenų apsauga. Tačiau nė viena iš šių nuostatų nėra tiesiogiai pritaikyta specifiniams algoritminio valdymo ir automatizuotos stebėsenos atvejams, kurie būdingi DI sistemoms. Žinoma tokiu atveju tikslinga vertinti ir DK 203 straipsnio apimtį ir DI naudojamų sistemų poveikį darbuotojo teisėms. Pagal Konstitucijos 22 straipsnį ir Europos žmogaus teisių konvencijos 8 straipsnį, darbuotojo privatumas saugomas net ir darbo vietoje. Šią poziciją aiškiai išplėtojo Europos Žmogaus Teisių Teismas (EŽTT) byloje *Barbulescu v. Romania* (Didžioji kolegija, 2017 m. rugsėjo 5 d., peticijos Nr. 61496/08). Teismas konstatavo, kad valstybės pozityvi pareiga saugoti asmens teisę į privatą gyvenimą apima ir darbo santykių kontekstą. Byloje sprendžiant dėl darbuotojo stebėjimo darbdavio iniciatyva, Teismas pabrėžė, kad bet kokia darbuotojo veiklos kontrolė, įskaitant elektroninių ryšių sekimą, turi būti aiškiai pagrįsta teisėtu tikslu, būtina demokratinėje visuomenėje, atitikti proporcingumo principą ir apie ją darbuotojas turi būti tinkamai, aiškiai ir iš anksto informuotas. Šios aplinkybės turi būti vertinamos atsižvelgiant į „subtilią pusiausvyrą“ tarp darbdavio interesų organizuoti veiklą ir darbuotojo teisės į privatą gyvenimą (*Barbulescu v. Romania*, 2017, §121–124). Ši jurisprudencija tampa vis reikšmingesnė algoritminės kontrolės ir emocinio stebėjimo priemonių kontekste, nes DI technologijos gali veikti automatiškai, be aiškaus žmogaus įsikišimo, ir analizuoti itin jautrius asmens duomenis – veido išraišką, emocijas, nuotaikų pokyčius ar net balso toną. Tokiose situacijose kyla rizika, kad darbuotojai apie šiuos procesus nebus tinkamai informuoti, o jų teisės į privatumą bus pažeistos. Todėl EŽTT byloje *Barbulescu v. Romania* išplėtoti kriterijai, tokie kaip informavimo pareiga, teisėtas tikslas ir priemonių proporcingumas, tampa ypač svarbūs vertinant naujus DI taikomus stebėsenos metodus darbo vietoje. Svarbią reikšmę biometrinių duomenų apsaugos kontekste turi Lietuvos vyriausiojo administracinio

teismo (LVAT) praktika. 2020 m. nutartyje teismas nagrinėjo situaciją, kai kelios įmonės, tarp jų UAB „Birštono šaltinis“, AB „Eglės sanatorija“ ir UAB „Namita“, naudojo darbuotojų pirštų atspaudus darbo laiko apskaitos ir drausmės kontrolei. Teismas konstatavo, kad tokia praktika pažeidė proporcingumo ir būtinybės principus, nes biometrinių duomenų rinkimas nebuvo nei būtinas, nei pagrįstas mažiau invazyvių priemonių nebuvimu. Be to, darbdaviai nesugebėjo įrodyti, jog pasirinktos priemonės atitinka BDAR 5 ir 9 straipsniuose įtvirtintus teisėtumo bei duomenų kiekio mažinimo principus. Teismas pabrėžė, kad darbdavys turi pareigą užtikrinti, jog asmens duomenų tvarkymas būtų ne tik teisėtas, bet ir proporcingas, o biometriniai duomenys – dėl savo ypatingo pobūdžio – gali būti tvarkomi tik išimtiniais atvejais. Ši byla parodė, kad net ir praktikoje paplitusios priemonės, tokios kaip pirštų atspaudų skenavimas, gali būti laikomos neteisėtomis, jei nėra tinkamai pagrįstos ir neatskleidžia duomenų subjekto teisių apsaugos priemonių.

EŽTT byla *Florindo de Almeida Vasconcelos Gramaxo v. Portugal* (2022) išryškina balansą tarp technologinės stebėsenos darbe ir darbuotojo teisės į privatą gyvenimą. Šioje byloje darbdavys naudojo GPS sistemą darbuotojo tarnybiniame automobilyje, o surinkti duomenys buvo panaudoti kaip įrodymas atleidimo procese. Teismas pripažino, kad nebuvo pažeistas Europos žmogaus teisių konvencijos 8 straipsnis, nes darbuotojas buvo aiškiai informuotas apie stebėseną, o jos tikslas ir apimtis buvo pagrįsti ir proporcingi (*Gramaxo v. Portugal*, 2022). Ši byla yra reikšminga vertinant DI pagrindu veikiančias stebėsenos sistemas, kurios dažnai renka vietas, elgsenos ar net biometrinius duomenis. Ji patvirtina, kad net ir išmaniosios stebėsenos priemonės gali būti teisėtos, jei tenkinamos sąlygos: aiškus informavimas, teisėtas tikslas ir būtinybė demokratinėje visuomenėje. DI sistemų diegimas turi atitikti šiuos kriterijus, ypač kai jos naudojamos darbuotojų judėjimo, veiklos ar rezultatų sekimui.

Pagal Bendrojo duomenų apsaugos reglamento (BDAR) 4 straipsnio 14 punktą, biometriniai duomenys yra „asmens duomenys, gauti naudojant konkrečius techninius apdorojimo metodus, susiję su fizinio asmens fizinėmis, fiziologinėmis ar elgesio savybėmis, leidžiančiomis arba patvirtinančiomis to fizinio asmens unikalų atpažinimą, pvz., veido atvaizdai ar pirštų atspaudai“ (Reglamentas (ES) 2016/679). Tokie duomenys darbo vietoje vis dažniau naudojami ne tik prisijungimui ir fizinės prieigos kontrolei, bet ir darbo našumo stebėsenai, elgesio analizei ar net emocinių reakcijų vertinimui naudojant DI sprendimus.

Darbo kodekso 26 straipsnyje įtvirtintas nediskriminavimo principas taip pat gali būti aktualus tuo atveju, jei biometrinės technologijos lemia nevienodą poveikį skirtingoms darbuotojų grupėms. Tokie skirtumai gali atsirasti dėl technologinio šališkumo (pvz., sistemų netikslumo tam tikrų etninių grupių atžvilgiu), sveikatos būklės (fizinių sutrikimų, psichologinių)

ir ar religinių įsitikinimų (kai biometrinis duomenų pateikimas prieštarautų ar nebūtų įmanomas dėl asmens religinių pažiūrų).

BDAR 9 straipsnyje biometriniai duomenys priskiriami prie specialiųjų kategorijų duomenų, kurių tvarkymas leidžiamas tik esant išimtims, pavyzdžiui, kai subjektas davė aiškų sutikimą, arba kai tvarkymas būtinas dėl darbo teisės normų ir laikantis tinkamų apsaugos priemonių (BDAR 9 str. 2 d.). Visgi, kaip pažymi Zaleskis (2019), darbo santykiuose toks sutikimas „dažnai būna fiktyvus – priklausomumo santykis daro jį formaliu ir ne laisvai išreikštu“ (p. 257–258). Darbo santykiuose egzistuoja pusiausvyros disbalansas – darbuotojas gali neturėti realios galimybės atsisakyti duoti sutikimą nepatirdamas pasekmių. Kaip pastebi Daria Bulgakova (2022), darbuotojo sutikimas neturi būti laikomas laisvu pasirinkimu, kai egzistuoja subordinacijos santykis su darbdaviu. Anot autorės, „tokiu atveju darbdavio ir darbuotojo santykis sukuria struktūrinį spaudimą, kuris de facto paneigia BDAR reikalaujamą sutikimo laisvumą“ (Bulgakova, 2022, p. 23–24). Šiuos principus patvirtina ir Europos duomenų apsaugos valdybos gairės 3/2019, kuriose teigiama, kad biometriniai duomenys neturėtų būti naudojami darbuotojų stebėsenai, nebent tai yra būtina, proporcinga ir pagrįsta konkrečiu darbo tikslu (EDPB, 2019, p. 11). Vokietijos, Italijos ir Nyderlandų duomenų apsaugos institucijų praktika, analizuota Bulgakovos tyrime, rodo, kad net ir darbuotojo sutikimas negali būti laikomas tinkama teisine baze. Autorė pabrėžia: „biometriniai duomenys prilyginami labai jautriems duomenims, kurių tvarkymas leidžiamas tik jei neįmanoma jų pakeisti mažiau invazyviu būdu, pvz., asmens pažymėjimu, slaptažodžiu ar žetonais“ (Bulgakova, 2022, p. 25–26). Todėl biometrinių duomenų naudojimas darbo aplinkoje negali būti grindžiamas vien sutikimu, būtinas proporcingumo, būtinybės ir informuotumo principų įgyvendinimas, kad nebūtų pažeista darbuotojo teisė į privatų gyvenimą ir duomenų apsaugą.

Dirbtinio intelekto veikimas neišvengiamai susijęs su duomenimis, kuo daugiau jų sistema gauna, tuo tiksliau ji geba modeliuoti sprendimus. Ši savybė tampa itin aktuali darbo santykiuose, kur automatizuoti sprendimai gali būti grindžiami tiek veiklos rezultatais, tiek biometriniais ar elgesio duomenimis. 2022 m. priimtas Duomenų valdymo aktas (Reglamentas (ES) 2022/868) nors ir nėra tiesiogiai skirtas darbo teisei, tačiau turi reikšmingą poveikį, kai darbdaviai naudojami trečiųjų šalių sprendimais, pavyzdžiui, paslaugomis, kurios apdoroja ar teikia duomenis sprendimų priėmimo sistemoms. Šiame akte nustatytos nuostatos formuoja teisinius rėmus, kurie reikalauja ne tik techninio atitikimo, bet ir vertybinių įsipareigojimų asmens teisių atžvilgiu. Viena iš reikšmingų akto nuostatų – reikalavimas užtikrinti skaidrumą. Duomenų tvarkymo logika, jų naudojimo paskirtis bei galimi duomenų gavėjai turi būti aiškiai ir suprantamai atskleisti, ypač kai sprendimai daro tiesioginę įtaką asmens profesinei padėčiai. Tai yra esminė sąlyga darbuotojui suvokti, kaip jo duomenys panaudojami atrankos, vertinimo ar

kitų svarbių sprendimų kontekste. Kitas svarbus principas yra nešališkumas. Aktas reikalauja, kad duomenų tarpininkai ir platformos veiktų taip, jog jų naudojamos sistemos neįtvirtintų sisteminių šališkumų ar nelygybių. Darbo santykiuose tai ypač reikšminga, nes bet koks algoritminis sprendimas dėl priėmimo į darbą, paaukštinimo ar atleidimo turi būti grindžiamas objektyviais kriterijais ir neturėti diskriminacinio poveikio. Galiausiai, duomenų tvarkymas turi būti pagrįstas sąžiningumu – ne tik teisiniu formalumu, bet ir realiu darbuotojų, kaip duomenų subjektų, interesų apsauga. Jei darbdavys naudoja sprendimų sistemomis, kurios grindžiamos trečiųjų šalių surinktais ar apdorotais duomenimis, jis išlieka atsakingas už tai, kad darbuotojas būtų tinkamai informuotas ir galėtų realizuoti savo teises, įskaitant prieigą prie duomenų, jų tikslumo patikrą ir sprendimo paaiškinimą. Tokiu būdu Duomenų valdymo aktas, net ir nebūdamas darbo teisės norminiu aktu, prisideda prie platesnio teisinio dirbtinio intelekto naudojimo konteksto kūrimo. Kartu su BDAR ir Dirbtinio intelekto aktu jis formuoja naują paradigmą, kurioje duomenys laikomi ne tik techniniu resursu, bet ir teisiškai reguliuojamu elementu, kuriam taikomi aiškūs atsakomybės, aiškumo ir žmogaus teisių apsaugos standartai.

Vienas iš aktualiausių ir dar retai praktikoje taikomų mechanizmų – poveikio duomenų apsaugai vertinimas (DPIA), numatytas Bendrojo duomenų apsaugos reglamento 35 straipsnyje. Šis vertinimas privalomas tais atvejais, kai duomenų tvarkymas gali sukelti didelę riziką duomenų subjekto (šiuo atveju – darbuotojo) teisėms ir laisvėms, pavyzdžiui, kai atliekama sisteminė stebėseną ar naudojamos naujos technologijos. Darbo teisės kontekste tai reiškia, kad darbdavys, planuodamas įdiegti DI sistemą, kuri vertins darbuotojų elgseną, produktyvumą ar net emocinę būklę, privalo atlikti išankstinį rizikos vertinimą ir dokumentuoti galimus pažeidimus. Europos duomenų apsaugos valdyba savo gairėse pabrėžia, kad DPIA nėra tik formalumas – tai praktinis įrankis, leidžiantis iš anksto įvertinti, ar planuojamos priemonės neprasilenkia su privatumo principais (EDPB, 2017).

Pagal BDAR 13 ir 14 straipsnius, duomenų valdytojas (darbdavys) privalo informuoti darbuotojus apie jų duomenų tvarkymą, įskaitant automatizuoto sprendimo logiką, tikslą ir galimas pasekmes. Vis dėlto, praktikoje gali susiklostyti situacijos, kai tokia informacija pateikiama itin bendro pobūdžio, formaliai, nesuteikiant darbuotojui realios galimybės suprasti, ar jo duomenys yra tvarkomi automatizuotai, kokia apimtimi ir ar sprendimas buvo priimtas be žmogaus įsikišimo. Tokie atvejai gali riboti BDAR 22 straipsnyje įtvirtintą teisę nebūti vien automatizuoto sprendimo objektu, ypač jei darbuotojui nesuteikiama galimybė pasinaudoti šio straipsnio 3 dalyje nurodytomis išimtimis – gauti žmogaus įsikišimą, išreikšti savo požiūrį, ginčyti sprendimą.

Lietuvos darbo kodekso 27 straipsnis šiuo atveju yra tiesiogiai taikytinas, nes įpareigoja darbdavį gerbti darbuotojo privatų gyvenimą ir užtikrinti asmens duomenų apsaugą. Šio

straipsnio 2 d. draudžia tvarkyti perteklinius duomenis, o 3 d. - reikalauja darbuotojus supažindinti su informacinių technologijų naudojimo ir stebėsenos darbo vietoje tvarka. Taigi biometrinių duomenų naudojimas gali būti laikomas teisėtu tik tuo atveju, jei jis būtinas darbo organizavimui, proporcingas tikslui ir nėra perteklinis.

Be teisėtumo, BDAR 5 straipsnyje įtvirtinti esminiai duomenų tvarkymo principai – teisėtumo, proporcingumo, būtinumo ir duomenų minimizavimo, kurie turi būti taikomi kiekvienu atveju. Daria Bulgakova (2022) atkreipia dėmesį, kad net ir egzistuojant teisei prievolei matuoti darbo laiką objektyviai bei patikimai, tai savaime nesuteikia pagrindo rinkti specialiųjų kategorijų duomenis, jeigu neįrodytas proporcingumo principas. Autorės teigimu, „nors darbdaviai gali teisėtai turėti interesą registruoti darbo laiką, tai nereiškia, jog pasirinkta techninė priemonė – pirštų atspaudų sistema – yra būtinai proporcinga“ (Bulgakova, 2022, p. 24). Kaip pabrėžia Europos duomenų apsaugos valdybos gairės 3/2019, biometrinių duomenų naudojimas darbuotojų stebėsenai gali būti pateisinamas tik tuo atveju, jei tokie metodai yra technologiškai būtini, proporcingi tikslui ir neturi mažiau invazinių alternatyvų (EDPB, 2019, p. 9, 73 p.).

Tarptautinėje praktikoje reikšmingas yra Jungtinės Karalystės Informacijos komisaro tarnybos (Information Commissioner's Office, ICO) sprendimas, priimtas 2024 m. balandį, kuriuo bendrovei „Serco Leisure“ buvo nurodyta nutraukti biometrinių duomenų – veido atpažinimo ir pirštų atspaudų naudojimą darbuotojų darbo laiko apskaitai. Tyrimo metu nustatyta, kad daugiau nei 2 000 darbuotojų biometriniai duomenys buvo tvarkomi be tinkamo teisinio pagrindo ir nesuteikiant alternatyvių apskaitos priemonių, taip pažeidžiant teisėtumo, proporcingumo ir duomenų kiekio mažinimo principus. Šis sprendimas sulaukė plataus rezonanso, po jo kitos įmonės, įskaitant „Virgin Active“, taip pat atsisakė panašių technologijų taikymo. Sprendimas aiškiai pabrėžė, kad darbdaviai turi įrodyti, jog biometrinių duomenų naudojimas yra būtinas, o ne pasirinktas dėl patogumo ar kontrolės stiprinimo.

Veidų atpažinimo sistemų diegimas biuruose ar sandėliuose darbuotojų įėjimo ir išėjimo kontrolei. Tokios sistemos, nors ir technologiškai efektyvios, kelia reikšmingų žmogaus teisių klausimų – jos gali būti laikomos perteklinėmis, ypač jei darbdavys neišnagrinėja mažiau invazinių alternatyvų, tokių kaip kortelės, kodai ar mobilioji autentifikacija. Daria Bulgakova (2022) pažymi, kad net jei darbdavys turi teisėtą interesą užtikrinti darbo laiką ar prieigos kontrolę, tai savaime nesuteikia pagrindo rinkti specialiųjų kategorijų duomenis. Anot jos, „darbdavio ir darbuotojo santykis sukuria struktūrinį spaudimą, kuris *de facto* paneigia BDAR reikalaujamą sutikimo laisvumą“ (Bulgakova, 2022, p. 23–24).

Be to, autorė teigia, kad proporcingumo principas reikalauja ne tik teisėto tikslo, bet ir technologinės būtinybės – pirštų atspaudų ar veido atpažinimo technologijos gali būti

netinkamos, jei egzistuoja mažiau intervencinės alternatyvos. Tokios priemonės turi būti ne tik efektyvios, bet ir pagrįstos orumo, informuotumo ir kontrolės pasidalijimo principais, nes priešingu atveju jos kuria nepermatomą, darbo aplinką, kuri gali pažeisti ar neužtikrinti darbuotojų teisių apsaugos.

Įdomų kontrastą su Europos požiūriu į biometrinių duomenų apsaugą atskleidžia Jungtinių Amerikos Valstijų Ilinojaus valstijos praktika. Čia galioja Biometric Information Privacy Act (BIPA, 740 ILCS 14), kuris reikalauja išankstinio rašytinio darbuotojo sutikimo prieš pradėdant bet kokią biometrinių duomenų rinkimą. Vienas reikšmingiausių šio įstatymo taikymo precedentų – *Rosenbach v. Six Flags Entertainment Corp.* (2019 IL 123186), kurioje Ilinojaus Aukščiausiasis Teismas pripažino, kad net ir vien procedūrinis pažeidimas (pvz., informavimo stoka ar sutikimo nebuvimas) yra pakankamas pagrindas teisiniam reikalavimui, neįrodinėjant materialinės žalos. Tai pabrėžia, kad šiame kontekste biometrinė teisė veikia kaip autonominė subjekto teisė į gynybą, nepriklausoma nuo žalos egzistavimo – kitaip nei Europos Sąjungos teisėje, kur žalos elementas vis dar dažnai laikomas būtinu ieškinio pagrindu (plg. BDAR 82 str.).

Biometrinių duomenų neteisėtas rinkimas ar naudojimas gali būti pagrindas civilinei atsakomybei. Lietuvos civilinio kodekso 6.250 straipsnis numato, kad neturtinė žala atlyginama, kai pažeidžiama asmens teisė į privatų gyvenimą ar duomenų apsaugą. Šis principas taikytinas ir tais atvejais, kai darbdavys naudoja dirbtinio intelekto pagrindu veikiančias sistemas, renkančias ar apdorojančias biometrinius duomenis be aiškaus teisinio pagrindo arba tinkamo darbuotojo informavimo. Tokiais atvejais, jei dėl pažeidimo atsiranda asmeniui reikšmingas poveikis (pvz., stebėjimo stresas, diskriminacija, atleidimas), galėtų būti sprendžiama dėl civilinės atsakomybės taikymo.

Kaip pabrėžia AI Now Institute (2023), algoritminis valdymas, įskaitant biometrinių duomenų naudojimą, kelia reikšmingą pavojų darbuotojų autonomijai, ypač kai tokios technologijos diegiamos be realios galimybės jas suprasti ar ginčyti jų poveikį. Pasak T. Davulio (2017), darbuotojo teisė į privatumą glaudžiai susijusi su orumo principu, todėl informacinio savarankiškumo užtikrinimas tampa būtina sąlyga šiuolaikiniuose darbo santykiuose, ypač atsižvelgiant į technologijų keliamą riziką (Davulis, 2017, p. 119).

Tuo tarpu Kellogg, Valentine ir Christin (2020) atkreipia dėmesį, kad algoritmai vis dažniau perima sprendimus dėl darbo organizavimo, tačiau „šios technologijos dažnai veikia taip, kad darbuotojui tampa neįmanoma žinoti, kas, kodėl ir kaip jį vertina“ – tokia situacija sukuria naują atsakomybės modelį, kai sprendimų priėmimas tampa išskaidytas tarp sistemos kūrėjų, operatorių ir naudotojų. Tai vadinama atsakomybės fragmentacija – darbuotojui tampa

sunku identifikuoti, kas konkrečiai atsakingas už priimtą sprendimą, todėl apsunkinamas teisių gynimo priemonių taikymas.

Viena iš labiausiai ginčytinų DI technologijų sričių darbo aplinkoje yra emocijų atpažinimo sistemos. Tokios sistemos analizuoja veido išraiškas, balso intonacijas, mikrojudesius ar kalbos struktūrą, siekdamos nustatyti darbuotojo ar kandidato emocinę būklę, stresą, nerimą, motyvaciją ar net „tinkamumą“ konkrečioms pareigoms. Jos vis dažniau taikomos automatizuotuose darbo pokalbiuose, darbuotojų stebėsenos procesuose ar klientų aptarnavimo srityse. Visgi šios priemonės yra vertinamos kritiškai tiek akademiniame, tiek reguliaciniame aplinkoje. Europos duomenų apsaugos valdyba (EDPB) ir Europos duomenų apsaugos priežiūros pareigūnas (EDPS) 2021 m. paskelbė bendrą pareiškimą, kuriame siūlė uždrausti emocijų atpažinimą viešosiose erdvėse, pabrėždami grėsmę žmogaus orumui ir duomenų apsaugai. Be to, šių sistemų taikymas gali pažeisti BDAR 9 straipsnyje įtvirtintus reikalavimus dėl specialiųjų kategorijų duomenų tvarkymo.

Europos duomenų apsaugos valdyba (EDPB) 2022 m. aiškiai pažymi, kad „emocijų atpažinimo sistemų naudojimas darbo vietoje beveik niekada negali būti laikomas proporcingu ir pagrįstu“, ypač atsižvelgiant į darbuotojo priklausomą padėtį ir sunkumus laisvai sutikti ar pasipriešinti tokiai stebėsenai (EDPB, 2022). Tokios sistemos laikomos itin invazyviomis, nes kišasi į darbuotojo vidinę būseną, kuri pagal žmogaus orumo principą neturėtų būti vertinama pagal objektyvius algoritminius kriterijus. Be to, kaip nurodoma Dirbtinio intelekto akto 5 straipsnio aiškinamuosiuose dokumentuose, šių sistemų veikimas dažnai grindžiamas stereotipinėmis prielaidomis apie „normalią“ emocinę išraišką, todėl jos gali būti kultūriškai šališkos ar net diskriminuojančios.

Dirbtinio intelekto aktas (Reglamentas (ES) 2024/1689) šią problematiką sprendžia tiesiogiai. 5 straipsnyje numatyta, kad emocijų atpažinimo ir biometrinio kategorizavimo sistemos viešose erdvėse, švietimo ir darbo aplinkoje, iš esmės draudžiamos, išskyrus labai ribotus išimčių atvejus. Šis reglamentinis požiūris iš esmės prilygsta šioms sistemoms taikomo principo inversijai – ne leidimui, o draudimui, nebent įrodoma aiški ir siaura būtinybė.

Nors emocijų atpažinimo technologijos dažnai kritikuojamos kaip pernelyg invazyvios, tam tikrose situacijose jų naudojimas gali būti laikomas teisėtu ir proporcingu. Pagal Dirbtinio intelekto akto 5 straipsnio 2 dalį, šios technologijos darbo aplinkoje gali būti taikomos tik išimtiniais atvejais – kai jos yra būtinos siekiant teisėto tikslo, pavyzdžiui, sveikatos ar saugos apsaugos, ir nėra kitų, mažiau įsikišimą reiškiančių priemonių (Reglamentas (ES) 2024/1689, 5 str. 2 d.). Vienas iš tokių pavyzdžių – emocijų analizės taikymas siekiant nustatyti perdegimo ar psichologinės įtampos riziką darbuotojams, dirbantiems ypatingo streso sąlygomis, kaip greitosios pagalbos medikams, ugniagesiams ar pilotams. Tokiu atveju, jei sistema yra skaidriai

įdiegta, dokumentuota ir tikrai būtina, ji gali būti suprantama kaip darbo saugos priemonė, o ne kontrolės instrumentas.

Tačiau net ir tokiais atvejais būtina užtikrinti, kad būtų laikomasi esminių duomenų apsaugos principų, įtvirtintų Bendrojo duomenų apsaugos reglamento (BDAR) 5 straipsnyje – teisėtumo, sąžiningumo, skaidrumo, tikslo apribojimo, duomenų kiekio mažinimo ir atskaitomybės. Tai reiškia, kad darbuotojas turi būti aiškiai ir suprantamai informuotas, kaip veikia emocijų analizės sistema, kokie duomenys renkami, kam jie naudojami, kiek laiko saugomi ir kokią įtaką gali turėti sprendimams dėl darbo. Be to, darbdavys turi įrodyti, kad tokia sistema tikrai būtina ir kad nėra kito, mažiau asmeniškai intervencinio būdo pasiekti tą patį tikslą. Šie reikalavimai grindžiami ne tik BDAR 5 ir 6 straipsniais, bet ir AI akte įtvirtintais darbuotojo apsaugos mechanizmais, taikomais aukštos rizikos DI sistemoms.

Tokia reguliavimo kryptis rodo, kad emocijų atpažinimas negali būti laikomas įprasta priemone darbuotojams vertinti ar stebėti. Atvirkščiai – tai yra jautri technologija, kurios naudojimas gali būti leidžiamas tik tuo atveju, jei darbdavys sugeba įrodyti jos būtinybę ir suderinti taikymą su darbuotojų pagrindinių teisių apsauga. Praktikoje tai reiškia, kad tokios sistemos turėtų būti naudojamos tik krašutiniais atvejais ir visada – su aiškiai apibrėžtomis apsaugos priemonėmis.

Reikšmingą nacionalinės praktikos akcentą pateikia advokatė Asta Macijauskienė (2024), analizuodama darbuotojams taikytinas taisykles dirbtinio intelekto įrankių naudojimui. Autorė pažymi, kad darbdaviai dažnai skatina DI naudojimą darbo vietoje neapibrėžę nei teisinio pagrindo, nei duomenų apsaugos ribų, nei atsakomybės už generuotą turinį. Tokia situacija kelia pavojų darbuotojo teisei į informavimą, ypač kai naudojant DI įrankius, priimti sprendimai gali daryti poveikį darbo kokybei ar rezultatams. Atsakomybės už dirbtinio intelekto sistemų generuojamą turinį klausimas darbo santykiuose metu išlieka neaiškus. Nors pagal bendruosius darbo teisės principus darbdavys atsako už savo vardu veikiančių sistemų, įskaitant technologijas, taikymą, DI naudojimo atveju gali kilti situacijų, kai darbuotojas priima sprendimus remdamasis sistemos pateikta informacija ar generuotais teksta, kurių turinys pažeidžia teisės normas (pvz., diskriminacija atrankos procese, klaidinanti informacija klientui). Tokiu atveju iškyla klausimas – ar atsako darbuotojas, kuris technologiškai nevaldė proceso, ar darbdavys, kuris sistemą įdiegė? Autorės Straipsnyje siūloma darbdaviams sukurti vidaus taisykles, kurios apibrėžtų leidžiamus įrankius, nustatytų atsakomybės paskirstymą bei užtikrintų darbuotojų autonomiją DI naudojimo kontekste. Dirbtinio intelekto akto 28 ir 29 straipsniai numato pareigą užtikrinti DI sprendimų paaiškinamumą ir galimybę priskirti veiksmus atsakingam subjektui. Tačiau nacionalinėje darbo teisėje tokio paskirstymo mechanizmo nėra – nėra aiškiai apibrėžta, kada darbuotojas gali būti atleistas nuo atsakomybės dėl DI padarytos

klaidos, arba kada darbdavys gali reikalauti atlyginti žalą, kilusią dėl DI generuoto turinio naudojimo.

Tai atitinka BDAR 13–15 straipsnių reikalavimus dėl aiškaus, skaidraus ir savalaikio informavimo apie duomenų tvarkymą bei automatizuotus sprendimus, tiek AI akto 29 straipsnio nuostatas, kurios įpareigoja darbdavius informuoti darbuotojus apie DI sistemų veikimą, paskirtį, sprendimų logiką bei galimą poveikį. Nacionalinėje teisėje šiuo metu nėra tiesiogiai įtvirtintos darbdavio pareigos dėl DI nustatyti tokias taisykles, tačiau jų papildymas ar sukūrimas praktikoje gali būti būtinas, siekiant įgyvendinti BDAR ir AI akto principus bei užtikrinti darbuotojų teisių apsaugą. Šiuo požiūriu, emocijų atpažinimo technologijos darbo santykiuose kelia ne tik teisinius, bet ir gilesnius žmogiškumo klausimus. Kai darbuotojo emocijos bandomos išmatuoti ir įvertinti algoritmais, kyla pavojus, kad vidiniai išgyvenimai bus suprasti paviršutiniškai ir neteisingai. Technologijos, kurios bando vertinti žmogaus emocijas, neturėtų pakeisti realaus žmogiško bendravimo ar empatijos. Todėl darbo teisės reguliavimas turi ne tik atitikti naujas technologines sąlygas, bet ir apsaugoti darbuotoją nuo pernelyg tolimos ar net neetiškos kontrolės, kuri gali pažeisti jo orumą ir psichologinį saugumą.

Šiame kontekste tampa itin reikšmingos ir praktinės gairės, kurios padeda darbdaviams orientotis tarp technologinio efektyvumo ir teisėtumo ribų. 2024 m. teisininkės Astos Macijauskienės suformuluotose rekomendacijose dėl dirbtinio intelekto taikymo darbo vietoje pabrėžiama ne tik duomenų apsaugos ir teisėtumo svarba, bet ir būtinybė užtikrinti darbuotojų informuotumą, aiškias duomenų tvarkymo ribas bei žmogaus įsikišimo garantiją. Kaip pažymi Macijauskienė (2024), „darbdavys privalo aiškiai apibrėžti, kokiems tikslams DI naudojamas, kokie duomenys renkami ir kaip jie analizuojami – nes skaidrumas nėra tik etinė, bet ir teisinė pareiga“. Tokie pasiūlymai rodo, kad technologinis progresas neturi eliminuoti žmogiškojo veiksnio, o priešingai, būti jam pavaldus. Tik taip galima išlaikyti darbo teisės esminį tikslą – apsaugoti žmogų, jo autonomiją ir teisę būti vertinamam ne tik kiekybinėmis analitinėmis priemonėmis, bet ir kaip asmenybę, veikiančią tarpusavio pasitikėjimu, pagarba ir lygiavertiškumu grindžiamuose santykiuose su darbdaviu.

2.2. Darbuotojų stebėseną, vertinimas, atrankų procesas

Tradiciškai darbo santykiai grindžiami tiesioginiu darbdavio ir darbuotojo bendravimu, individualiu konteksto vertinimu bei aiškiai apibrėžta atsakomybe už sprendimus. Tačiau DI pagrindu veikiančios sistemos keičia šią paradigmą, kai sprendimų dėl kandidatų atrankos, darbo krūvio paskirstymo, našumo vertinimo ar net darbo sutarties nutraukimo priėmimas vykdomas pasitelkiant DI algoritmus. Toks pokytis kelia ne tik duomenų apsaugos ar techninio reguliavimo iššūkius, bet ir fundamentalius klausimus dėl teisinės atskaitomybės, sprendimų pagrįstumo, darbuotojo teisėtų lūkesčių bei žmogaus teisių apsaugos. Dirbtinio intelekto technologijų taikymas darbuotojų atrankoje laikomas vienu iš reikšmingiausių žmogiškųjų išteklių valdymo skaitmenizavimo etapų. Pastaraisiais metais organizacijos vis dažniau pasitelkia automatizuotus sprendimų priėmimo mechanizmus, kurie leidžia apdoroti didelius kandidatų duomenų kiekius, analizuoti CV turinį, prognozuoti kandidato tinkamumą konkrečioms pareigoms ar net analizuoti neverbalinę komunikaciją nuotolinių interviu metu.

Dirbtinio intelekto integravimas į darbo santykius ypač pastebimas darbuotojų atrankos procese, kur vis dažniau taikomi automatizuoti sprendimų modeliai: algoritmai, analizuojantys gyvenimo aprašymus, elgsenos modelius ar net emocinę raišką pokalbių metu. Šios sistemos gali būti grindžiamos istoriniais duomenimis, todėl kyla rizika, kad jos atkartos diskriminacines praktikas, pagrįstas lytimi, amžiumi, etnine kilmė ar kitais draudžiamais pagrindais. Kaip pabrėžia De Stefano (2020), algoritminė atranka keičia patį žmogaus vertinimo principą – nuo individualios analizės pereinama prie statistinių tikimybių ir koreliacijų, kurios gali būti socialiai šališkos. Europos Sąjungos dirbtinio intelekto akto (2024/1689) 6 ir 34 straipsniuose darbuotojų atrankos sistemoms priskiriamas „didelės rizikos“ statusas, todėl darbdaviai, naudojančios tokias sistemas, privalo užtikrinti jų skaidrumą, žmogaus priežiūrą ir paaiškinamumą. Tai reiškia, kad algoritmas negali būti vienintelis veiksnys, nulėmęs sprendimą dėl kandidato priėmimo ar atmetimo, būtinas žmogaus įsikišimas, kaip numato ir BDAR 22 straipsnis. Taigi, darbuotojo teisinė padėtis šiuo atveju reikalauja papildomų garantijų – ne tik formalaus duomenų apsaugos užtikrinimo, bet ir realaus, skaidraus, proporcingo ir individualiai pagrįsto sprendimo.

Pavyzdžiui, kaip pažymi Ineta Breskienė, algoritminis valdymas remiasi tuo, kad „DI sistemos automatiškai analizuoja surinktus duomenis ir formuoja tam tikras išvadas, kurios gali turėti tiesiogines teises pasekmes“ (Breskienė, 2024, p. 104). Toks metodas, pritaikytas darbuotojų atrankos kontekste, leidžia efektyviau atlikti išankstinį kandidatų filtravimą. Vis dėlto, šioje situacijoje išryškėja keli fundamentaliai svarbūs iššūkiai: algoritminio šališkumo (angl. *algorithmic bias*), skaidrumo ir duomenų apsaugos problematika.

Europos Parlamento ir Tarybos reglamentas (ES) 2024/1689 (Dirbtinio intelekto aktas) aiškiai nustato, kad darbo kontekste taikomos didelės rizikos DI sistemos, įskaitant sistemas, naudojamas vertinant kandidatus ar priimant sprendimus dėl jų įdarbinimo, privalo atitikti griežtus skaidrumo, duomenų valdymo ir žmogaus priežiūros reikalavimus (Dirbtinio intelekto aktas, 2024, 6 str. 1 d., a p.).

Reikšmingu pavyzdžiu, atskleidžiančiu dirbtinio intelekto technologijų taikymo darbuotojų atrankoje pavojus, laikytina „Amazon“ korporacijos įgyvendinta, bet galiausiai nutraukta personalo atrankos sistema, pagrįsta DI sprendimų priėmimu. 2014 m. bendrovė pradėjo kurti algoritminę atrankos sistemą, kuri automatizuotai vertino kandidatų gyvenimo aprašymus, suteikdama jiems įvertinimus nuo vienos iki penkių žvaigždučių. Tačiau netrukus paaiškėjo, kad sistema ėmė nuosekliai mažinti moterų kandidatūrų vertinimus. Tai lėmė mokymo duomenų bazė, pagrįsta istoriniais CV, kuriuose dominavo vyrai, o frazės, susijusios su moteriška veikla (pvz., „moterų šachmatų klubas“), buvo ignoruojamos ar traktuojamos kaip neigiami požymiai. Šis atvejis atskleidžia, kaip netinkamai suformuoti treniravimo duomenys gali įtvirtinti struktūrinį šališkumą ir pažeisti lygiateisiškumo principą (Dastin, 2018). Ši situacija, nors ir nesukėlusi formalios teismo bylos, konceptualiai atspindi galimą 1964 m. JAV Civilinių teisių akto pažeidimą, kuris aiškiai draudžia bet kokią diskriminaciją įdarbinimo procese dėl lyties (Title VII, 42 U.S.C. § 2000e-2). „Amazon“ atvejis leidžia konstatuoti, kad algoritminės atrankos sistemos, kai jos apmokomos šališkais duomenimis ir veikia be žmogaus peržiūros, gali sukurti sistemines netiesioginės diskriminacijos rizikas. Tokie atvejai atskleidžia būtinybę subalansuoti technologinį efektyvumą su teisiniu įsipareigojimu užtikrinti nediskriminavimo principą – tiek nacionaliniu, tiek tarptautiniu lygiu. Europos Sąjungos pagrindinių teisių chartijos 21 straipsnyje tiesiogiai įtvirtinta draudimo diskriminuoti nuostata grindžia ir AI akte išdėstytus reikalavimus dėl algoritminių sprendimų skaidrumo, paaiškinamumo ir kontrolės. Tai suponuoja, kad DI taikymas įdarbinimo kontekste reikalauja ne tik technologinio tobulumo, bet ir žmogaus teisių apsaugos prioriteto. Priešingu atveju, tai gali sukelti dideles pasekmes darbuotojų teisių apsaugai.

Vienu pirmųjų reikšmingų teismo sprendimų, įteisinančių galimybę taikyti atsakomybę už algoritminę diskriminaciją darbuotojų atrankoje, laikytina byla *Mobley v. Workday Inc.* (3:23-cv-00770), nagrinėta JAV Šiaurės Kalifornijos apygardos teisme. 2024 m. liepos 12 d. teisėja Rita F. Lin priėmė išvadą, kad ieškovas Derekas Mobley nurodė pakankamai faktinių aplinkybių, leidžiančių pagrįstai manyti, jog Workday Inc., teikusi automatizuotus įdarbinimo sprendimų įrankius, galėjo diskriminuoti jį dėl rasės (afroamerikietis), amžiaus (virš 40 metų) ir negalios (nerimo ir depresijos sutrikimai), pažeidžiant JAV federalinius teisės aktus – *Civil Rights Act*, *ADEA* ir *ADA*. Teismas taip pat pažymėjo, kad algoritminis sprendimų priėmimo mechanizmas,

veikiantis kaip „darbdavio agentas“, galėjo atlikti diskriminacines funkcijas, nepriklausomai nuo tyčios buvimo. Šis sprendimas reikšmingai išplečia atsakomybės taikymą DI kūrėjams bei diegėjams, kai jų algoritmai turi sisteminių neigiamą poveikį pažeidžiamoms visuomenės grupėms.

Bylos *Mobley v. Workday Inc.* kontekste reikšminga tampa faktinė aplinkybė, kad ieškovas pateikė daugiau nei 100 darbo paraiškų bendrovėms, kurios naudojo Workday algoritmus kandidatams vertinti. Visos paraiškos buvo atmestos, nesuteikus galimybės sužinoti sprendimo motyvų. Teismas akcentavo, kad automatizuota sistema veikė kaip pirmasis ir iš esmės lemiantis atrankos proceso etapas, eliminuodamas kandidatus be žmogaus peržiūros. Kaip nurodyta sprendime, „ieškovas įvardija pakankamai aplinkybių, iš kurių galima spręsti, kad Workday algoritmai atliko funkcijas, priskirtinas darbdaviui, įskaitant kandidatų atmetimą“ (*Mobley v. Workday Inc.*, 3:23-cv-00770, Dkt. 46, p. 9). Tokia pozicija atveria kelią DI sistemų tiekėjų teisinei atsakomybei, kai jų algoritmai tampa ne tik technine priemone, bet realiu sprendimų priėmėju, darančiu tiesioginį poveikį asmens galimybei dalyvauti darbo rinkoje.

Nors teismas atmetė kaltinimus dėl tyčinės diskriminacijos motyvo, dėl nepakankamo įrodymų pagrindo apie diskriminacinį ketinimą, jis aiškiai pripažino galimą „netiesioginį diskriminacinį poveikį“ (angl. *disparate impact*), kurį sukelia automatinės sistemos, veikiančios remiantis istoriniais, socialiai šališkais duomenimis. Tai atvėrė kelią tęsti bylą ir iš esmės įteisino galimybę DI sistemų kūrėjus laikyti atsakingais pagal federalinius antidiskriminacinius teisės aktus, jei jų sukurti algoritmai veikia kaip atrankos filtrai ir sukelia sistemine žalą tam tikroms socialinėms grupėms. Precedento prasme ši byla laikytina esminiu posūkiu, kuris signalizuoja, kad automatizuotos sprendimų sistemos nebegali būti traktuojamos kaip neutrali priemonė, jos tampa teisiškai atsakingos už sprendimų padarinius, ypač kai jų veiksmai paveikia pažeidžiamas grupes. Tokiu būdu formuojama jurisprudencija, kurioje technologijų teikėjai įtraukiami į atsakomybės sferą dėl galimų diskriminacinių padarinių darbo teisės srityje.

Svarbu pažymėti, kad byloje nebuvo išsamiai atskleista algoritminė sistema ar techninė jos struktūra – šis aspektas lieka ateities bylos nagrinėjimo klausimu. Tačiau teismas akcentavo, kad net ir nesant žinių apie konkrečią technologinę diskriminacijos mechaniką, DI tiekėjas negali išvengti atsakomybės vien dėl to, kad sistema yra „juodoji dėžė“ – nepermatoma ar techniškai sudėtinga. Tokiu būdu byla atspindi iššūkį šiuolaikinėje darbo teisėje: algoritminės atrankos sistemų sprendimai grindžiami sudėtingomis statistinėmis prognozėmis, kurių veikimo logikos negali paaiškinti nei darbuotojas, nei darbdavys. Teisinėje doktrinoje tai kelia klausimą dėl įrodinėjimo naštos paskirstymo ir teisės į paaiškinimą, kaip tai numatyta ir pagal BDAR 22 straipsnį ES teisėje. Teismas savo sprendimu taip pat iškėlė precedentą, pagal kurį trečiosios šalies technologijų tiekėjas, kaip DI sprendimo kūrėjas ir diegėjas, gali būti laikomas „darbdavio

agentu“ pagal antidiskriminacinius įstatymus, jei algoritmas *de facto* lemia sprendimą dėl įdarbinimo. Šis požiūris teisiškai patvirtina, kad algoritminis sprendimas, net ir neturėdamas tyčinio diskriminacinio ketinimo, gali būti neteisėtas, jeigu jo rezultatai sistemingai neigiamai paveikia tam tikras visuomenės grupes. Todėl darbdaviai privalo užtikrinti žmogaus įsikišimą į sprendimų priėmimą.

DI sistemų diegimas darbuotojų veiklos kontrolėje iš esmės transformuoja tradicinius darbo santykių pagrindus, pereinant nuo individualaus darbdavio sprendimo prie algoritminio valdymo (angl. *algorithmic management*) modelio. Toks valdymas grindžiamas nuolatinio darbuotojų duomenų rinkimu, jų analizavimu realiuoju laiku ir automatizuotu sprendimų priėmimu dėl darbo grafikų, užduočių paskirstymo, produktyvumo vertinimo ar net drausminių veiksmų inicijavimo. Ineta Breskienė šiuos procesus kvalifikuoja kaip „valdymo formą, kurioje sprendimai apie darbuotojo elgseną ar likimą priimami algoritmų pagrindu, remiantis iš jų išvestomis įžvalgomis“, ir pažymi, kad tokia kontrolės forma kelia didelę grėsmę darbuotojo teisėms, ypač privatumo, informavimo, saviraiškos ir orumo srityse (Breskienė, 2024, p. 104).

Šiuos aspektus tiesiogiai reglamentuoja 2024 m. priimtas Europos dirbtinio intelekto aktas (Reglamentas (ES) 2024/1689), kuris priskiria DI sistemas, naudojamas darbuotojų stebėsenai, vertinimui ir elgsenos prognozavimui, prie didelės rizikos sistemų (6 str. 1 d. b p.). Tai leidžia daryti išvadą, kad tokios sistemos leidžiamos tik laikantis griežtų reikalavimų dėl skaidrumo, žmogaus priežiūros, poveikio vertinimo, duomenų tikslumo ir nešališkumo. Svarbu nepamiršti, kad pagal ES Bendrąjį duomenų apsaugos reglamentą (BDAR), darbdaviai privalo informuoti darbuotojus apie automatizuoto sprendimų priėmimo taikymą (13–14 str.), o jei sprendimas gali sukelti reikšmingas pasekmes – darbuotojas turi teisę į paaiškinimą bei sprendimo peržiūrą žmogaus (22 str.).

Svarbiu jurisprudenciniu precedentu laikytina 2023 m. byla *EEOC v. iTutorGroup, Inc.* (Civil Action No. 2:22-cv-07182), kurioje Jungtinių Valstijų Lygių galimybių įdarbinimo komisija (EEOC) pateikė pirmąją federalinę ieškinį dėl diskriminacijos, kylančios iš dirbtinio intelekto pagrindu veikiančios įdarbinimo sistemos. EEOC teigimu, anglų kalbos pamokas teikianti įmonė „iTutorGroup“, naudodama automatizuotą kandidatų vertinimo algoritmą, automatiškai atmesdavo paraiškas vien remdamasi pareiškėjų gimimo metais – eliminuodama moteris vyresnes nei 55 metų ir vyrus vyresnius nei 60 metų. Sistema neveikė nei statistiniu pagrindu, nei dėl netyčinio šališkumo – diskriminacija buvo įdiegta programiškai, tiesiogiai pažeidžiant 1967 m. Age Discrimination in Employment Act (ADEA), draudžiantį naudoti amžių kaip įdarbinimo kriterijų. Byla baigėsi taikos sutartimi, pagal kurią bendrovė įsipareigojo sumokėti 365 000 dolerių kompensaciją nukentėjusiems, bei peržiūrėti algoritmų veikimą ir įgyvendinti kontrolės mechanizmus. Šis atvejis laikytinas jurisprudenciniu lūžiu, nes jis įtvirtino

atsakomybės principą tiek technologijų naudotojams, tiek jų kūrėjams – net ir tais atvejais, kai diskriminacija nėra sąmoninga, o kyla iš pačios algoritminės logikos. Jungtinių Valstijų Lygių galimybių įdarbinimo komisijos (EEOC) vadovė Charlotte A. Burrows akcentavo, kad „technologijų naudojimas neatleidžia darbdavių nuo atsakomybės laikytis antidiskriminacinių teisės aktų“, taip nubrėždama principinę ribą tarp technologinio efektyvumo ir teisinio įpareigojimo užtikrinti žmogaus teises. Ši pozicija stiprina ne tik JAV doktriną, bet ir tarptautinę teisės plėtrą, įskaitant ES Dirbtinio intelekto aktą (Reglamentas (ES) 2024/1689), kuriame darbuotojų atrankos ir vertinimo sritys tiesiogiai priskiriamos didelės rizikos DI taikymo sritims (6 str. 1 d. b p.). Ši byla rodo, kad bet koks automatizuotas sprendimų priėmimas darbo teisėje reikalauja ne tik technologinio tikslumo, bet ir sisteminės atsakomybės už pasekmes, galinčias pažeisti diskriminacijos draudimo principą bei teisę į vienodą prieigą prie darbo galimybių.

Europos Sąjungos pagrindinių teisių chartijos 30 straipsnis įtvirtina darbuotojo teisę į apsaugą nuo nepagrįsto atleidimo, grindžiamą teisėtumo, proporcingumo ir socialinio dialogo principais. Šią teisę išplečia ir Europos Žmogaus Teisių Teismo praktika – sprendime *Lombardi Vallauri v. Italy* (2009) akcentuota, kad bet kokio neigiamo sprendimo atžvilgiu asmuo turi teisę būti iš anksto informuotas, išklašytas ir dalyvauti priimant sprendimą, kuris paveikia jo profesinį statusą. Tačiau algoritminio atleidimo atveju šie procesiniai principai dažnai nėra įgyvendinami – darbuotojas neturi galimybės nei suprasti algoritmo logikos, nei apskūsti sprendimo realiam asmeniui, tik darbdaviui. Tai akivaizdžiai prieštarauja Bendrojo duomenų apsaugos reglamento (BDAR) 22 straipsniui, kuris draudžia vien automatizuotą sprendimų priėmimą, jei jis sukelia teisines pasekmes, ir įpareigoja darbdavį suteikti galimybę sprendimą peržiūrėti žmogui bei paaiškinti jo loginę struktūrą. Esant tokiame sprendimui, darbuotojas netenka galimybės suprasti, koku pagrindu buvo įvertinta jo. Tokia padėtis silpnina darbuotojo procesines garantijas ir tai, kad darbdavio sprendimai būtų grindžiami galiojančiais teisės aktais, pagrįsti ir nuspėjami. Reikšmingu jurisprudenciniu precedentu laikytina 2021 m. byla *Riders Union v. Uber & Ola* (ECLI:NL:RBAMS:2021:5305), kurioje Amsterdamo apygardos teismas pripažino, kad platforminiai darbuotojai buvo atleisti vadovaujantis vien algoritminiu sprendimu dėl tariamo „įtartino elgesio“. Teismas nustatė, kad šis sprendimas buvo priimtas be žmogaus įsikišimo, pažeidžiant pagrindinius teisingo proceso reikalavimus – darbuotojai nebuvo informuoti apie sprendimo logiką, neturėjo galimybės pateikti prieštaravimų ar užginčyti sprendimą. Teismas konstatavo, kad algoritmas veikė pažeisdamas Bendrojo duomenų apsaugos reglamento (BDAR) 22 straipsnio nuostatas, kurios draudžia sprendimus su reikšmingomis pasekmėmis, kai jie priimami vien automatizuotai. Šis atvejis laikytinas reikšmingu, signalizuojančiu, kad automatizuoti sprendimai darbo teisėje negali būti laikomi teisėtais, jei

neįgyvendintos procedūrinės garantijos – teisė į paaiškinimą, žmogaus įsikišimą bei ginčijimo galimybę.

Atsižvelgiant į tai, dirbtiniu intelektu pagrįstas darbo sutarties nutraukimas turi būti vertinamas per kelis teisinius filtrus: (a) ar buvo informuotas darbuotojas, kad sprendimas priimamas automatizuotai; (b) ar sprendimas gali būti paaiškintas loginiais principais; (c) ar suteikta galimybė žmogaus intervencijai; (d) ar sprendimo pasekmės atitinka proporcingumo kriterijus. Šie kriterijai kyla iš Bendrojo duomenų apsaugos reglamento (BDAR) 22 straipsnio, taip pat iš nacionalinio darbo teisės reguliavimo – Lietuvos darbo kodekso 57 ir 58 straipsnių. Kaip pažymi Juškevičiūtė-Vilienė (2020), dirbtinio intelekto sprendimų automatizavimas, ypač kai kalbama apie nuotolinius ar automatinius algoritminius sprendimus, gali pažeisti ne tik teisę į lygybę ir privatų gyvenimą, bet ir fundamentalią teisę į veiksmingą teisminę gynybą (Juškevičiūtė-Vilienė, 2020, p. 126–127). Šios rizikos tampa itin aktualios darbo teisės kontekste, kai algoritminiai sprendimai gali lemti įdarbinimą, darbo nutraukimą ar vertinimą, tačiau darbuotojui nesuteikiama galimybė suprasti sprendimo pagrindo. Lieka tik galimybė ginčyti patį sprendimą. Tokiais atvejais gali būti pažeidžiamos ne tik Europos Sąjungos pagrindinių teisių chartijos 21 straipsnyje įtvirtintas nediskriminavimo principas, bet ir 47 straipsnyje numatyta teisė į veiksmingą teisinę gynybą. Europos Komisija savo 2020 m. Baltojoje knygoje dėl dirbtinio intelekto (COM(2020) 65 final) akcentavo, kad didelės rizikos srityse, tarp jų ir darbuotojų atrankoje bei vertinime, būtina užtikrinti skaidrumą, atskaitomybę bei žmogaus kontrolę, kad automatizuoti sprendimai nepažeistų žmogaus teisinių garantijų. Tokie reikalavimai atliepia ir Lietuvos Respublikos Konstitucinio Teismo jurisprudenciją. 2014 m. balandžio 16 d. nutarime teismas aiškiai pabrėžė, kad teisė į teisingumą negali būti vien deklaratyvi – ji turi būti realiai įgyvendinama, užtikrinant kiekvienam asmeniui galimybę veiksmingai ginčyti sprendimą, kuris pažeidžia jo teises ar teisėtus interesus (Konstitucinis Teismas, 2014). Ši doktrina galėtų būti taikoma ir dirbtinio intelekto pagrindu priimamiems sprendimams, kurie sukelia teises pasekmes darbuotojui.

Kai darbdaviai naudoja automatizuotas sistemas darbuotojų vertinimui, iškyla svarbių klausimų dėl to, kiek darbuotojas turi realios įtakos su juo susijusiems sprendimams ir kaip tai keičia įprastą darbdavio ir darbuotojo santykį. Dirbtiniu intelektu pagrįstos sistemos perkelia sprendimų galią iš žmogaus darbdavio į technologinę sistemą, kuri dažnai veikia netransparentiškai ir be aiškos atsakomybės. Tokia technologinė tvarka dehumanizuoja darbo santykius, sprendimai dėl darbuotojo vertinimo ar atleidimo priimami be jo supratimo ar galimybės ginčyti pačią logiką, kai ji nežinoma. Tai pažeidžia pagrindinius teisės principus – teisėtumo, proporcingumo bei teisėtų lūkesčių apsaugos ir kelia klausimą, ar tokie algoritminiai mechanizmai gali būti laikomi tinkamu pagrindu reikšmingiems darbo sprendimams. Todėl vis

plačiau pripažįstama, kad algoritminė kontrolė turi būti vertinama kaip tokio pat lygmens teisinio reguliavimo objektas, kaip ir bet kuris kitas darbdavio sprendimas, galintis sukelti teises pasekmes darbuotojui (ETUI, 2023).

Europos duomenų apsaugos valdyba (EDPB) savo 2020 m. gairėse pabrėžia, kad pagal BDAR 22 straipsnį asmuo turi teisę nesutikti su sprendimu, kuris yra priimtas vien automatizuotai ir sukelia teises arba reikšmingas pasekmes. Toks sprendimas gali būti teisėtas tik tuomet, kai užtikrinamos papildomos garantijos – žmogaus įsikišimas, teisė pareikšti nuomonę ir ginčyti sprendimą. Šios nuostatos suponuoja pareigą taikyti ne tik technines, bet ir teises bei procedūrinės apsaugos priemones, kurios užtikrintų skaidrų ir pagrįstą sprendimų priėmimo procesą. Lietuvos Respublikos darbo kodekso 26 straipsnyje tiesiogiai įtvirtinta darbdavio pareiga taikyti vienodus atrankos, vertinimo ir atleidimo kriterijus visiems darbuotojams, neatsižvelgiant į jų amžių, lytį, kilmę, sveikatos būklę ar kitas savybes, kurios negali būti teisėtu pagrindu riboti dalyvavimą darbo rinkoje. Net ir algoritminiais sprendimais grįstos atrankos ar vertinimo procedūros turi būti vertinamos pagal tas pačias teisės normas, kurios taikomos visiems darbdavio veiksams, galintiems turėti teisinės reikšmės darbuotojo padėčiai. Tik toks vertinamas galėtų būti dar griežtesnis.

Lietuvos darbo teisėje, ypač taikant darbo sutarties nutraukimo pagrindus pagal Darbo kodekso 57 straipsnį – kai darbdavys nutraukia sutartį dėl svarbių priežasčių – numatyta pareiga sprendimą pagrįsti ir pateikti darbuotojui rašytinį motyvavimą. Darbdavys turi įrodyti, kad sutarties nutraukimas atitinka teisėtumo, proporcingumo ir nediskriminavimo principus. Be to, DK 64 straipsnis įtvirtina formalų nutraukimo įforminimo standartą, kuris suponuoja darbuotojo teisę suprasti jam taikomo sprendimo pobūdį ir pasekmes. Šios nacionalinės nuostatos turi būti aiškinamos kartu su BDAR 22 straipsniu ir ES Dirbtinio intelekto akto 29 straipsniu, kurie draudžia vien automatizuotą sprendimų priėmimą, turintį reikšmingą poveikį asmeniui. Todėl, jei sprendimas dėl darbo sutarties nutraukimo grindžiamas dirbtinio intelekto atliktu darbuotojo vertinimu, darbuotojui turi būti suteikta galimybė reikalauti žmogaus įsikišimo, gauti sprendimo paaiškinimą ir jį apskųsti. Šiuo metu Lietuvos teisėje tokios situacijos aiškiai nereglamentuotos, todėl nacionalinės normos turėtų būti papildytos tam, kad užtikrintų skaidrumą, sprendimo prieinamumą ir darbuotojo procesines teises, kai nutraukimo priežastys kyla iš DI veikimo logikos.

Kaip pastebi Davulis (2008), darbo teisės modernizavimas reikalauja ne vien formaliai pritaikyti senąsias normas prie naujų ekonominių ar technologinių aplinkybių, bet ir užtikrinti, kad tokios transformacijos nepažeistų darbuotojo teisinės padėties, informacinio savarankiškumo bei darbo teisės socialinės funkcijos, kuri nukreipta į darbuotojo kaip silpnesnės šalies apsaugą. . Autoriaus teigimu, darbo teisė negali atsisakyti savo socialinės funkcijos net ir globalizacijos ar

technologinių inovacijų akivaizdoje. Tai ypač aktualu atleidimo iš darbo situacijose, kurios yra glaudžiai susijusios su žmogaus saugumu, pasitikėjimu sistema ir teise į sąžiningą procesą.

Todėl Lietuvos darbo teisės sistemoje kyla būtinybė aiškiau reglamentuoti, koku mastu darbdavys gali priimti darbo sutarties nutraukimo sprendimą, pagrįstą vien dirbtinio intelekto sistemos rekomendacija ar automatizuotu modeliu. Nors galiojantis Darbo kodeksas nereglamentuoja šios situacijos tiesiogiai, tokio pobūdžio sprendimai, nesant žmogaus įsikišimo, gali kelti abejonių dėl atitikties esminiams darbo teisės principams – teisėtumo, proporcingumo, skaidrumo ir darbuotojo teisės būti informuotam. Aiškinant DK nuostatas sistemiskai kartu su BDAR 22 straipsniu bei Europos Parlamento ir Tarybos reglamento (ES) 2024/1689 29 straipsniu, galima daryti išvadą, kad darbuotojui turėtų būti užtikrinta teisė žinoti sprendimo pagrindus, gauti paaiškinimą. Nors šiuo metu Darbo kodeksas nereikalauja, kad sprendimą būtinai priimtų fizinis asmuo, o ne DI sistema, atsižvelgiant į konstitucinius ir Europos Sąjungos teisės standartus, tokio pobūdžio sprendimai turėtų būti peržiūrimi ir patvirtinami žmogaus, turinčio diskreciją ir atsakomybę. Priešingu atveju, kyla pavojus, kad darbuotojo teisės nebus užtikrintos procesiniu lygmeniu, nors formaliai ir būtų laikomasi DK nustatytų procedūrų.

Nors tarptautiniu lygmeniu vis daugiau dėmesio skiriama dirbtinio intelekto sistemų taikymui darbo teisės santykiuose, Lietuvos teisinėje praktikoje ši sritis dar nėra reflektuota per precedentų ar teisminės analizės prizmę. Nacionalinėje jurisprudencijoje iki šiol nėra nagrinėtų bylų, kuriose būtų sprendžiami ginčai dėl darbuotojų teisių pažeidimų, kilusių dėl algoritminio valdymo ar automatizuotų sprendimų taikymo darbo atrankose ar vertinimo, atleidimo, priėmimo procesuose. Siekiant užkirsti kelią esamų normų interpretacijai, tikslinga būtų atlikti reglamentavimo papildymus konkrečiai dėl DI. Tai, kad Lietuvos Respublikos darbo kodeksas nenumato specifinio reglamentavimo algoritminiam vertinimui ar automatizuotiems sprendimams, ypač atrankos ar stebėsenos kontekste, reiškia, kad šie reiškiniai turi būti vertinami taikant bendrąsias darbo teisės ir asmens duomenų apsaugos normas, įskaitant proporcingumo, teisėtumo, skaidrumo bei darbuotojo informavimo principus (Reglamentas (ES) 2016/679; Europos Parlamento ir Tarybos reglamentas (ES) 2024/1689).

Pažymėtina, kad Reglamento (ES) 2024/1689 29 straipsnyje įtvirtinta pareiga užtikrinti žmogaus įsitraukimą ir priežiūrą tais atvejais, kai dirbtinio intelekto sistemos naudojamos didelės rizikos situacijose, taikytina ir darbo teisinių santykių kontekste, kuomet sprendimai daro tiesioginį poveikį asmens teisei padėčiai. Šiuo atžvilgiu itin svarbus tampa Bendrojo duomenų apsaugos reglamento 22 straipsnis. Šių nuostatų integravimas į nacionalinę praktiką rekomenduotina sąlyga siekiant išvengti struktūrinės darbuotojų teisių apsaugos spragų.

Darytina prielaida, kad teismų vaidmuo ateityje bus itin reikšmingas šios srities doktrinos plėtrai – tiek aiškinant taikytinus principus, tiek formuojant jurisprudenciją, kurioje dirbtinio

intelekto technologijų taikymas būtų vertinamas proporcingai darbuotojo pagrindinių teisių apsaugos poreikiams. Tokiu būdu būtų formuojama integruota pozicija, derinanti technologinę pažangą su socialinės apsaugos tikslais, kaip jie apibrėžiami tiek nacionalinėje, tiek Europos žmogaus teisių konvencijos ir Europos socialinės chartijos kontekste.

2.3. Automatizuotų sprendimų poveikis

Kandidatai gali būti neinformuojami, kad jų paraiškų vertinimui taikomos automatizuotos sprendimų priėmimo sistemos. Tokia praktika gali sukelti abejonių dėl atitikties Bendrojo duomenų apsaugos reglamento (Reglamentas (ES) 2016/679) 13 ir 22 straipsniams. BDAR 22 straipsnis numato, kad: „Duomenų subjektas turi teisę nebūti sprendimo, kuris yra grindžiamas tik automatizuotu duomenų tvarkymu subjektu, jeigu toks sprendimas sukelia jam teisinės pasekmės arba reikšmingai paveikia.“ Tokiais atvejais automatizuoti sprendimai, jei jie priimami be žmogaus įsikišimo, aiškaus sutikimo arba nesant specialaus teisinio pagrindo, gali būti laikomi nesuderinamais su BDAR reikalavimais. Europos duomenų apsaugos valdyba savo gairėse (EDPB, 2020, p. 21) pabrėžia, kad sprendimai, darantys reikšmingą poveikį – įskaitant įdarbinimą – negali būti priimami vien automatizuotai, jei nėra žmogaus įsikišimo ir papildomų garantijų. Šią poziciją detalizuoja ir Zaleskis (2019), nurodydamas, kad BDAR ne tik riboja tokių sprendimų taikymą, bet ir suponuoja pareigą užtikrinti jų paaiškinamumą bei žmogaus dalyvavimą vertinimo procese. Nors BDAR nedetalizuoja išreikštos teisės į algoritminio sprendimo logikos atskleidimą, ši teisė, Zaleskio nuomone, gali būti sistemiškai išvedama iš informavimo, skaidrumo ir sąžiningumo principų (Zaleskis, 2019, p. 257).

Darbo teisės kontekste automatizuotas sprendimas dėl atleidimo ar nepriėmimo gali kelti abejonių dėl atitikties teisiniams reikalavimams, jei jis priimamas be žmogaus įsikišimo, o darbuotojui ar kandidatui nesuteikiama galimybė suprasti jo pagrindimo ar jį užginčyti. Tokiais atvejais duomenų apsaugos teisė tampa funkciškai susijusi su darbo teise, o Bendrasis duomenų apsaugos reglamentas (BDAR) veikia ne tik kaip techninių duomenų tvarkymo standartų rinkinys, bet ir kaip darbuotojo teisių apsaugos šaltinis. BDAR 22 straipsnyje įtvirtinta, kad asmuo turi teisę nebūti vien automatizuoto sprendimo, turinčio teisinių pasekmių, subjektu, išskyrus tam tikras išimtis, ir kad turi būti užtikrinta galimybė gauti paaiškinimą bei įsikišimą žmogui. Nors Lietuvos Respublikos darbo kodeksas šiuo metu tiesiogiai nenumato kandidato teisės reikalauti paaiškinimo dėl nepriėmimo sprendimo, ypač kai jis pagrįstas automatizuota analize, BDAR šią pareigą įtvirtina, ypač kai sprendimas turi reikšmingų pasekmių asmeniui. Tačiau reali tokios teisės įgyvendinimo galimybė dažnai priklauso nuo to, kiek darbdavys geba atskleisti DI sistemos veikimo logiką. Kai sprendimai priimami remiantis algoritmais, kurių veikimo principai yra nepermatomi, vadinamosiomis „juodosios dėžės“ sistemomis, teisė gauti paaiškinimą formaliai egzistuoja, tačiau praktiškai gali būti sunkiai realizuojama.

Dirbtinio intelekto technologijų diegimas lemia struktūrinius pokyčius darbo rinkoje. Nors technologinis progresas tradiciškai siejamas su produktyvumo augimu, šiuolaikiniai procesai sukelia sisteminę socialinę riziką, ypač kai sprendimai dėl darbo funkcijų

optimizavimo priimami be darbuotojų įtraukimo ar socialinio dialogo. Kaip pažymi Bitinas, Tartilas ir Litvaitienė (2011), socialinės apsaugos teisėje vienas iš esminių principų yra asmens apsauga nuo nepalankių ir neprognozuojamų socialinių pokyčių, darančių poveikį jo ekonominiam saugumui. Todėl tokios situacijos, kai darbuotojas netenka darbo ar socialinių garantijų dėl automatizuoto sprendimo, turi būti vertinamos kaip socialinės rizikos atvejai, kuriems būtina taikyti atitinkamas teisinės apsaugos priemones – pavyzdžiui, užtikrinti teisę į išėtinę išmoką, perkvalifikavimą, ar išplėsti darbuotojo teises dalyvauti sprendimų priėmimo procese.

Europos Komisijos 2021 m. pateiktame Platforminio darbo direktyvos projekte (COM(2021) 762 final) išskiriama būtinybė užtikrinti žmogaus įsitraukimą į visus darbo santykiams reikšmingus sprendimus, ypač kai jie grindžiami automatizuotais procesais. Direktyvos 6 straipsnis įpareigoja darbdavius informuoti darbuotojus apie algoritmų poveikį jų darbo rezultatams, darbo grafikui ir sutarties tęstinumui. Tai rodo, kad Europos Sąjungos teisėje sprendimai, galintys lemti darbo praradimą, turi būti pagrįsti ne tik efektyvumo, bet ir skaidrumo, žmogaus kontrolės bei socialinės apsaugos principais.

Tarptautinės darbo organizacijos (ILO) ataskaitoje nurodoma, kad iki 2030 m. apie 14 % darbo vietų pasaulyje gali būti visiškai automatizuotos, o dar 32 % – iš dalies (ILO, 2017, p. 1). Šios tendencijos ypač aktualios vyresnio amžiaus, mažesnės kvalifikacijos darbuotojams bei regioninėse darbo rinkose, kuriose dominuoja lengvai automatizuojami procesai. De Stefano ir Wouters (2022) pažymi, kad „nors DI suteikia galimybių darbo našumui, jis taip pat rizikuoja perkelti didžiąją socialinę kainą ant darbuotojų pečių, kurie negali prisitaikyti pakankamai greitai“ (p. 12). Ši įžvalga atskleidžia, kad technologinis progresas turi būti vertinamas ne tik ekonominio efektyvumo, bet ir socialinio teisingumo požiūriu.

Atsižvelgiant į jų poveikį darbuotojo teisėms, dirbtinio intelekto priemonių naudojimas darbo santykiuose neturėtų būti laikomas vien techniniu ar neutraliai organizaciniu sprendimu. Tokios sistemos gali turėti tiesioginių teisinių pasekmių, todėl jų taikymas turėtų būti aiškiai reglamentuotas, įvertinant ne tik efektyvumą, bet ir socialinį poveikį bei žmogaus teisių užtikrinimą. Nacionaliniame teisiniame reguliavime reikalinga įtvirtinti apsaugos priemones, kurios leistų darbuotojams iš anksto žinoti apie DI sistemų naudojimą, suprasti jų veikimo principus, reikalauti žmogaus peržiūros tais atvejais, kai sprendimai lemia jų teisinę ar profesinę padėtį.

Automatizuotų sprendimų diegimas darbo teisėje iškelia sudėtingą klausimą dėl sprendimų kilmės ir atsakomybės paskirstymo. Tradiciškai sprendimai dėl darbo organizavimo, veiklos vertinimo ar net sutarties nutraukimo priklausė darbdavio diskrecijai, su kuria buvo siejama ir atsakomybė. Tačiau kai šie sprendimai naudaujant DI sistemas yra formalizuojami

darbdavio vardu, bet faktiškai yra generuojami algoritminių sistemų, atsiranda teisinio atirbojimo problema: darbuotojui tampa sudėtinga nustatyti, kas iš tiesų priėmė sprendimą, ir koku mastu jis gali būti ginčijamas. Žinoma, darbuotojo atžvilgiu už tokių sprendimų priėmimą vis tiek lieka atsakingas darbdavys.

Atsižvelgiant į Bendrojo duomenų apsaugos reglamento ir Dirbtinio intelekto akto reikalavimus, automatizuoti sprendimai darbo santykiuose negali būti laikomi neutraliais ar savaime teisėtais. Jie privalo būti įtraukti į darbo organizavimo struktūrą taip, kad būtų užtikrinta darbuotojo teisė būti informuotam apie sprendimo priėmimo logiką, galimybė pareikšti savo nuomonę ir reikalauti žmogaus įsikišimo. Kaip pažymi Cefaliello ir Kullmann (2022), net jei algoritminės sistemos techniškai atitinka galiojančius teisės aktus, jų „socialinis teisėtumas“ priklauso nuo to, ar darbuotojas realiai geba suprasti sprendimų logiką ir turi galimybę ją įvertinti bei ginčyti (p. 174). Toks požiūris padeda atkurti pusiausvyrą tarp technologinės sprendimų galios ir darbuotojo dalyvavimo teisės, stiprinant pasitikėjimą sprendimų priėmimo procesu.

Siekiant užtikrinti darbuotojo teisių apsaugą algoritminio valdymo kontekste, darbdaviai turėtų darbo vidaus taisyklėse įtvirtinti aiškias nuostatas dėl DI sistemų veikimo principų, sprendimų priėmimo hierarchijos ir informavimo tvarkos. Tai apimtų informavimo mechanizmą apie automatizuotą vertinimą, darbuotojo teisę susipažinti su sprendimų formavimo logika (bent apytikrių modelių), bei prievolę peržiūrėti reikšmingus automatizuotus sprendimus vadovybės lygiu.

Lietuvos Respublikos darbo kodekso 206 straipsnis aiškiai reglamentuoja darbdavio pareigą konsultuotis su darbo taryba tvirtinant vietinius norminius teisės aktus, kurie turi tiesioginę įtaką darbuotojų darbo sąlygoms, privatumo apsaugai, darbo stebėsenai ar informacinių technologijų naudojimui. Vis dėlto ši norma taikoma tik tais atvejais, kai yra patvirtinami ar keičiami konkretūs teisės aktai (pvz., vidaus taisyklės, politikos dokumentai). Tokiu būdu, jei darbdavys diegia dirbtinio intelekto (DI) sistemą be formalizuoto norminio akto, pavyzdžiui, integruodamas automatinį darbo našumo vertinimo įrankį ar personalo atrankos algoritmą per išorinę platformą – DK 206 straipsnio dispozicija neapima tokio sprendimo formaliai.

Darbo kodekso 203 straipsnyje įtvirtinta darbuotojų ir jų atstovų teisė į informavimą ir konsultavimą darbdavio sprendimų, susijusių su darbuotojų darbo, socialinių, ekonominių teisių bei interesų įgyvendinimu ir gynimu, klausimais. Ši formuluotė iš pirmo žvilgsnio apima itin platų situacijų spektrą, suteikdama norminį pagrindą konsultacijoms ir dėl naujų technologijų įdiegimo, jeigu jos turi įtakos minėtoms darbuotojų teisėms. Atsižvelgiant į tai, kad dirbtinio intelekto sprendimai gali lemti esminius pokyčius darbo pobūdyje, pareigų paskirstyme, veiklos

vertinime, netgi tapti netiesioginiu pagrindu arba pagalbininku dėl sprendimo darbo sutarties nutraukimo atžvilgiu, akivaizdu, kad šių sistemų diegimas darbdavio iniciatyva turi tiesioginį ryšį su darbuotojo teisėmis ir interesais.

Vis dėlto šio straipsnio taikymas DI kontekste šiuo metu galėtų palikti vietos interpretacijai, kadangi DI, algoritminis valdymas ar automatizuoti sprendimai normoje tiesiogiai neįvardijami. Tai padidintų tikimybę, kad tiek darbdaviai, tiek darbo tarybos gali nevienodai suprasti pareigos konsultuotis apimtį, ypač jei sprendimas dėl DI diegimo, naudojimo priimamas be vietinio norminio akto, o techninės naujovės įvedamos per išorinių paslaugų platformas ar kitas integruotas sistemas. Toks neapibrėžtumas ypač aktualus sparčiai diegiant išorinius algoritminio vertinimo įrankius ar darbuotojų veiklos stebėjimo sistemas, kurios daro įtaką darbuotojams darbo santykiuose, bet nėra formaliai įtvirtinamos vietiniais norminiais aktais (kas aktyvuotų DK 206 straipsnio mechanizmą).

Šioje vietoje verta palyginti Lietuvos situaciją su Vokietijos teisine raida. 2021 m. birželio 18 d. Vokietijoje priimtas Betriebsrätmodernisierungsgesetz – Įmonių tarybų modernizavimo įstatymas, kuris papildė Vokietijos Įmonių tarybų įstatymą (Betriebsverfassungsgesetz, BetrVG) specifinėmis nuostatomis apie dirbtinio intelekto naudojimą darbo vietoje. Šis teisės aktas buvo parengtas kaip atsakas į darbo pasaulio skaitmenizaciją ir algoritminės kontrolės stiprėjimą, siekiant stiprinti darbuotojų atstovų galias šių procesų kontrolėje. DI tematika įtraukta tiesiogiai į kelis esminius BetrVG straipsnius, taip nepaliekant teisinės abejonės dėl darbuotojų tarybų kompetencijos ir darbdavių pareigų.

Pavyzdžiui, BetrVG § 90 (1) 3 p. numato darbdavio pareigą laiku ir išsamiai informuoti darbuotojų tarybą apie planuojamą DI naudojimą darbo procesuose. § 87 (1) 6 p. nustato bendro sprendimo teisę tarybai dėl bet kokių sistemų, kuriomis galima stebėti darbuotojų elgesį ar našumą. Šis punktas apima DI sistemas, kurios vykdo analizę ar vertinimą. § 95 (2a) aiškiai reglamentuoja, kad DI naudojimas personalo atrankos gairėse privalo būti derinamas su taryba. Galiausiai, § 80 (3) įtvirtina darbuotojų tarybos teisę pasitelkti išorinį ekspertą DI klausimams vertinti darbdavio lėšomis, ir šis poreikis laikomas preziumuojamu, nebereikalaujant papildomo pagrindimo.

Vokietijos reguliavimo modelis formuotas būtent siekiant užtikrinti, kad technologinės inovacijos darbo vietoje nepažeistų darbuotojų pagrindinių teisių. Nors pačiame įstatyme sąvoka „dirbtinis intelektas“ (vok. *künstliche Intelligenz*) nėra apibrėžta, aiškinamajame įstatymo rašte pateikiamas plačiai suprantamas DI kontekstas. DI įvardijamas kaip, nedeterministinės, adaptyvios sistemos, generuojančios neprognozuojamus rezultatus. Toks sąvokos aiškinimas leidžia taikyti nuostatas ne tik siaurai techninėms programoms, bet ir platesniam DI sistemų spektrui.

Atsižvelgiant į tai, Lietuvos teisėje iki šiol nėra aiškaus atitikmens, kuris įtvirtintų privalomą konsultavimąsi su darbuotojų atstovais konkrečiai DI diegimo ar naudojimo atveju. Nors, kaip analizuota, DK 203 straipsnio konstrukcija teoriškai apima tokias situacijas, jos taikymas gali būti grindžiamas vien interpretacija, svarstant, ar DI įdiegimas turės esminį poveikį darbuotojo teisėms, jei galutinius sprendimus vis tiek priima žmogus, o toks interpretavimas neužtikrina dabartinio DK vienodo taikymo ir darbuotojų teisių užtikrinimo.

Todėl būtų tikslinga svarstyti galimybę Lietuvos darbo kodeksą papildyti aiškiais nuostatomis, įtvirtinančiomis pareigą konsultuotis su darbuotojų atstovais dėl DI sistemų diegimo, nepriklausomai nuo to, ar tvirtinamas vietinis norminis aktas. Tokios nuostatos galėtų aiškiai atriboti, kada privaloma konsultacija, DI diegiamas per išorines platformas. Taip pat sustiprintų darbuotojų atstovų institucines galias prižiūrėti DI sprendimų įgyvendinimą ir atitiktų ES DI akto (2024/1689) ir BDAR formuojamą požiūrį į darbuotojų informavimą apie DI.

Algoritminis valdymas darbo aplinkoje apima ne tik sprendimų priėmimą dėl įdarbinimo ar atleidimo, bet ir darbuotojų veiklos planavimą, užduočių paskirstymą bei našumo kontrolę, pagrįstą nuotoliniu duomenų sekimu. Todoli-Signes (2021) pastebi, kad tokios sistemos dažnai naudojamos vertinant darbuotojų produktyvumą ar priimant sprendimus dėl darbo santykių, net kai darbuotojas nežino nei vertinimo logikos, nei atsakomybės paskirstymo tarp darbdavio ir algoritmo tiekėjo. Empiriniai tyrimai rodo, kad tokio pobūdžio automatizuotas valdymas, ypač kai sprendimų kriterijai išlieka neaiškūs, gali kelti padidėjusį darbuotojų stresą, neapibrėžtumo jausmą ir pasitikėjimo mažėjimą. Pasak ILO, algoritminis vertinimas dažnai vykdomas be aiškių kriterijų ir realaus žmogaus įsikišimo, todėl tokia kontrolė gali lemti nuolatinį stresą, neapibrėžtumo jausmą bei spaudimą darbuotojui (ILO, 2022, p. 14–15).

Šie algoritminio valdymo mechanizmai daro poveikį ne tik darbuotojo teisinėms garantijoms, bet ir emocinei būklei. Europos profesinių sąjungų konfederacijos (ETUC) analizė pabrėžia, kad intensyvi algoritminė kontrolė, kai nėra aiškiai apibrėžtų taisyklių, žmogaus priežiūros ar derybinių garantijų, gali prisidėti prie emocinio išsekimo, motyvacijos mažėjimo ir net darbo funkcijų atsisakymo, ypač kai darbuotojas susiduria su nenuosekliu, nuasmenintu vertinimu (ETUC, 2023). Tokie poveikiai susiję ne tik su psichosocialine našta, bet ir su struktūriniu darbuotojo socialinio balso praradimu, kai algoritminė kontrolė iš esmės apeina kolektyvinio atstovavimo ir dialogo mechanizmus. Panašiai, Rimkutė (2022, p. 89) pažymi, kad automatizuotos sprendimų sistemos „formuoja objektyvios kaltės priskyrimo sumaištį“, o tai sukuria ne tik teisinio neapibrėžtumo, bet ir socialinio nesaugumo lauką darbuotojui. Tokia „sumaištis būseną“ tampa nuolatinio darbuotojo palydovu skaitmenizuotoje darbo aplinkoje, kur jo kontrolės galimybės ribotos, o atsakomybės paskirstymas – fragmentuotas. Visa tai kelia klausimą, ar algoritminis valdymas gali būti suderinamas su darbuotojo psichologiniu saugumu

ir darbuotojo galimybe suprasti bei ginčyti sprendimus, darančius poveikį jo padėčiai darbo santykiuose.

Atsižvelgiant į tai, akivaizdu, kad automatizuotų sprendimų naudojimas darbo santykiuose turi būti vertinamas ne vien per duomenų apsaugos ar skaidrumo prizmę, bet ir psichosocialinio poveikio aspektu. Taip pat, atkreiptinas dėmesys į būtinybę darbo tvarkos taisyklėse ir DK reglamentuoti DI taikymo aspektus darbo santykiuose, kurie apimtų taikymo ribas, informavimą, atsakomybę. Šie aspektai šiuo metu nėra aiškiai reglamentuoti Lietuvos darbo kodekse, būtent DI taikymo atvejais, o tai sudaro prielaidas tiek teisiniam, tiek darbo teisių pažeidimams.

3. ATSAKOMYBĖS PASKIRSTYMAS UŽ DIRBTINIO INTELEKTO NAUDOJIMĄ DARBO TEISINIUOSE SANTYKIUOSE

3.1. Darbdavių, darbuotojų ir trečiųjų šalių atsakomybės aspektai

Dirbtinio intelekto sistemų diegimas darbo santykiuose iš esmės transformuoja klasikinius atsakomybės paskirstymo principus tarp darbdavio ir darbuotojo. Tradiciškai sprendimų dėl darbuotojo priėmimo, vertinimo ar atleidimo priėmimas buvo neatsiejamai susijęs su darbdavio diskrecijos teise ir atitinkama atsakomybe. Tačiau algoritminio valdymo kontekste šie sprendimai dažnai formaliai priskiriami darbdaviui, bet faktiškai inicijuojami arba determinuojami DI sistemos, kurios veikimo logika gali būti neprieinama, nepermatoma arba kintanti dėl mokymosi procesų. Kyla klausimas dėl teisinės atsakomybės lokalizavimo: ar darbdavys gali būti laikomas vieninteliu atsakingu subjektu, kai sprendimo logika slypi algoritme? Ši problema ypač aktuali, kai DI sistemą sukūrė arba administruoja trečioji šalis – tiekėjas ar platforma, neturintys tiesioginio ryšio su darbo teisinais subjektais. Lietuvos teisės nacionaliniai aktai šiuo metu nėra visapusiškai pritaikyti tokioms situacijoms, jis remiasi žmogaus sprendimo prezumpcija ir aiškiai nenumato technologinių sprendimų ribų, atsakomybės paskirstymo algoritminio sprendimo atveju.

ES Dirbtinio intelekto akte (Reglamentas (ES) 2024/1689) skiriami du pagrindiniai DI sprendimų grandinės subjektai – „teikėjas“ (angl. *provider*) ir „naudotojas“ (*user*). Teikėjui priskiriama atsakomybė už sistemos architektūrą, dokumentaciją, rizikų įvertinimą ir saugą, o naudotojui – už praktinį sprendimų įgyvendinimą ir žmogaus priežiūrą. Aukštos rizikos DI sistemų naudotojai, kaip tai apibrėžta 16–22 straipsniuose, privalo užtikrinti sprendimų skaidrumą, galimybę peržiūrėti sprendimą bei jo atsekamumą.

Darbo teisės kontekste svarbu aptarti ne tik darbuotojų teises dėl DI naudojimo, bet ir darbdavio galimybę kreiptis į kūrėją, tiekėją. Lietuvos nacionalinėje teisėje šiuo metu neegzistuoja aiškiai suformuluotas teisinis režimas, kuris leistų darbdaviui pareikšti regresinį reikalavimą dirbtinio intelekto sistemos tiekėjui ar kūrėjui tuo atveju, kai tokios sistemos veikimas gali sukelti neteisėtas pasekmes – pavyzdžiui, klaidingą sprendimą dėl darbuotojo atleidimo ar diskriminacinę elgesį. Civilinės teisės sistemoje tokio pobūdžio atsakomybė teoriškai galėtų būti grindžiama Civilinio kodekso 6.263 straipsnyje įtvirtintu bendroju deliktinės atsakomybės principu – žalos padarymo pagrindu, taip pat 6.270 straipsniu, reglamentuojančiu gamintojo atsakomybę už nesaugų produktą. Tačiau nei viena iš šių normų kol kas nėra tiesiogiai pritaikyta algoritminėms sistemoms, kurios veikia skaitmeniniu, dinamišku ir dažnai nepermatomu principu. Dėl šios priežasties ypač sudėtinga nustatyti priežastinį ryšį ar atriboti atsakomybę tarp sistemos kūrėjo, tiekėjo ir naudotojo. Toks

reguliacinis neapibrėžtumas lemia, kad darbdavys išlieka vienintelis atsakingas subjektas, net jei neturėjo realios galimybės kontroliuoti ar numatyti sprendimo turinio. Ši situacija atskleidžia poreikį sistemiškai peržiūrėti atsakomybės mechanizmus ir svarstyti galimybę įtvirtinti horizontalų, aiškų rizikos paskirstymo bei prevencinės atsakomybės modelį, pritaikytą dirbtinio intelekto specifikai.

Atsižvelgiant į šiuo metu egzistuojančius teisinio reguliavimo trūkumus nacionalinėje teisėje, ypač regresinio reikalavimo kontekste, svarbiu instrumentu tampa Europos Komisijos pasiūlytas Dirbtinio intelekto atsakomybės direktyvos projektas (AILD), kuriuo siekiama suvienodinti civilinės atsakomybės taisykles už žalą, padarytą naudojant DI sistemas visoje Europos Sąjungoje. Vienas iš šios direktyvos tikslų – užtikrinti, kad nukentėjusieji galėtų efektyviai pareikšti pretenzijas dėl žalos, kilusios dėl DI sistemos veikimo, nepriklausomai nuo to, ar tiesioginis valdytojas buvo darbdavys, ar trečiasis subjektas, pavyzdžiui, sistemos kūrėjas ar tiekėjas. Direktyvos preambulėje akcentuojama, kad nacionaliniuose teisiniuose režimuose dažnai nėra priemonių, leidžiančių darbuotojui ar net darbdaviui reikalauti atlyginti žalą, kai algoritminė sistema elgiasi nenumatytai arba diskriminuojančiai tiesiogiai iš kūrėjo, tiekėjo. AILD siūlo įtvirtinti atgręžtinės atsakomybės (angl. *recourse liability*) galimybę, leidžiančią darbdaviui, kaip pirminei atsakomybei tenkančiam subjektui, reikšti reikalavimą technologijų tiekėjui arba modeliavimo paslaugos teikėjui, jeigu nustatoma, kad pastarasis nesilaikė priežiūros ar rizikos mažinimo pareigų. Tai leistų išskaidyti riziką tarp visų algoritminės ekosistemos dalyvių bei sudarytų prielaidas sukurti solidarumo principu grindžiamą teisinės atsakomybės sistemą. Tačiau būtina pažymėti, kad kol AILD nėra galutinai priimta ir inkorporuota į Lietuvos teisinę sistemą, regresinio reikalavimo galimybės nacionalinėje teisėje išlieka neapibrėžtos, fragmentiškos ir priklauso nuo bendrųjų deliktinės bei sutartinės atsakomybės taisyklių taikymo.

Tiekėjui galima atsakomybė už DI sistemos architektūrinę kokybę – pradinę algoritminę logiką, mokymo duomenų rinkinius, modelio veikimo prognozuojamumą ir dokumentaciją. Naudotojas, šiuo atveju darbdavys, atsako už tai, kaip sistema pritaikoma darbo santykių kontekste: kaip vertinami darbuotojai, ar užtikrinamas žmogaus įsikišimas ir ar sprendimai atitinka teisinio pagrįstumo reikalavimus. Ši atsakomybės sandūra sukuria kompleksinę teisinę situaciją, kai algoritminiai sprendimai sukelia pasekmes, tačiau neaišku, ar už jas turi atsakyti tik darbdavys, ar taip pat DI kūrėjas. Nors AI aktas įtvirtina naudotojo pareigas užtikrinti sprendimų skaidrumą ir žmogaus priežiūrą (Reglamentas (ES) 2024/1689, 16–22 str.), atsakomybės atribojimo mechanizmai nėra detaliai išplėtoti ir paliekami nacionalinių teisės sistemų reguliavimui.

Europos Komisijos parengtas DI atsakomybės direktyvos pasiūlymas (2022) pabrėžia, kad „kai DI sistema laikoma aukštos rizikos, jos naudotojui – dažniausiai darbdaviui – tenka

teisinė pareiga įrodyti, jog buvo imtasi visų būtinų atsargumo priemonių“ (Europos Komisija, 2022, p. 8). Taigi, net jeigu klaidą padaro DI sistema, teisinė atsakomybė vis tiek tenka darbdaviui, jeigu jis nesugebėjo įrodyti tinkamos sistemos priežiūros ir atitikimo teisės aktams.

Kaip pažymi Čerka, Grigienė ir Sirbikytė (2015), klasikinės deliktinės atsakomybės konstrukcijos tampa ribotos, kai jos taikomos dirbtinio intelekto kontekste, ypač tais atvejais, kai DI sistemos veikia autonomiškai, o jų sprendimai negali būti aiškiai susieti su konkretaus subjekto valia, kontrole ar nuspėjamumu. Tokiose situacijose esminiai deliktinės teisės kriterijai – kaltė, numatomumas ir galimybė užkirsti kelią žalai – tampa sunkiai taikomi, o tai destabilizuoja atsakomybės paskirstymo logiką. Reaguodami į šį conceptualų iššūkį, autoriai siūlo funkcinį atsakomybės modelį, pagal kurį DI kūrėjas atsako už sistemos architektūrą, algoritminės logikos ribotumus ir informavimo pareigą, o naudotojas – už tai, kaip ši sistema veikia konkrečiu socialiniu ir teisiniu kontekstu. Taip perkeliama atsakomybės ašis nuo veikimo priežasties prie sprendimų struktūrinės inžinerijos, taip užtikrinant proporcingesnę rizikų paskirstymą ir teisės taikymo adaptaciją prie technologinių transformacijų.

Kaip pažymi Bartkus ir Stundys (2018), algoritminis sprendimų priėmimas nediskvalifikuoja darbdavio iš teisės ir moralės perspektyvos – priklausomybė nuo DI sistemos nesumažina jo pareigos pagrįsti sprendimą, užtikrinti jo teisėtumą ir paisyti darbuotojo orumo darbo santykiuose. Net jeigu sprendimas yra formaliai automatizuotas, darbdavys privalo užtikrinti, kad jis būtų žmogui suprantamas, atitiktų teisinį skaidrumą bei būtų pagrįstas aiškiais kriterijais. Tokia pozicija iš esmės reiškia, kad DI sistemų taikymas neišlaisvina darbdavio nuo tradicinės atsakomybės pareigos ypač kai algoritminis sprendimas turi tiesioginių pasekmių darbuotojo teisėms ar statusui. Bartkus ir Stundys iškelia esminį principą: automatizavimas yra techninis sprendimas, bet atsakomybė – visada teisinė. Tai reikalauja iš darbdavio ne tik techninio budrumo, bet ir sisteminio gebėjimo užtikrinti, kad sprendimo pasekmės atitiktų darbuotojų teisių apsaugos principus, kaip jie apibrėžiami tiek nacionaliniuose teisės aktuose, tiek Europos Sąjungos teisėje.

Probleminė situacija kyla tuomet, kai darbdavys, priimdamas sprendimą dėl darbuotojo, remiasi DI sistemos pateikta analize ar vertinimo rezultatais, tačiau neturi aiškios galimybės suprasti sistemos veikimo logikos ar įvertinti jos tikslumo. Tokiais atvejais darbuotojo teisė į sprendimo pagrindimą, ginčijimą ar žmogaus įsikišimą gali būti ribota, o darbdaviui kyla rizika prisiimti atsakomybę už sprendimą, kurio pagrindai nėra iki galo aiškūs ar patikrinti. Manytina, kad visais atvejais turi būti užtikrinamas žmogaus įsikišimas priimant tokius sprendimus, taip ši rizika galėtų būti labiau suvaldoma. Minėtu atveju klasikiniai civilinės atsakomybės elementai – neteisėti veiksmai, kaltė, priežastinis ryšys ir žala – tampa sunkiai pritaikomi, nes sprendimas nėra tiesiogiai išvestas iš žmogaus valios. Kaip pažymi Čerka, Grigienė ir Sirbikytė (2015, p.

377–378), tokiose situacijose būtina taikyti diferencijuotą požiūrį: kūrėjas turėtų atsakyti už algoritminės sistemos architektūrą, ribas ir logikos skaidrumą, o naudotojas, už sistemos taikymą konkrečiame teisiniame kontekste bei atsakomybės neužtikrinimą. Tokiu būdu pagrindžiamas vadinamasis funkcinės atsakomybės modelis, kuris darbo teisėje gali būti taikomas atibojant kūrėjo ir naudotojo atsakomybę: vienas atsako už konstrukcinį defektą, kitas – už netinkamą eksploatavimą. Šis galėtų būti aktualus sprendžiant ginčus, kai darbuotojas neturi galimybės suprasti ar užginčyti sprendimą, nors jis grindžiamas DI sistema.

Šiuolaikinės dirbtinio intelekto sistemos, veikiančios mokymosi principais, kelia esminius iššūkius deliktinės atsakomybės doktrinai. Atsakomybė grindžiama tuo, ar asmuo galėjo numatyti žalą ir jos išvengti. Tačiau mokymosi pagrindu veikiantys algoritmai dažnai veikia pagal tikimybinę ir adaptyvią logiką, kuri gali būti nenuspėjama tiek kūrėjui, tiek naudotojui. Dėl šios priežasties akademiniėje doktrinoje siūlomas perėjimas nuo rezultato kontrolės prie vadinamojo projekto valdomumo modelio, kuriame atsakomybė siejama su tuo, ar sistema buvo tinkamai suprojektuota, testuota, dokumentuota ir ar naudotojas buvo tinkamai informuotas apie jos ribotumus (Wachter ir Mittelstadt, 2019)

Mokslinėje doktrinoje vis dažniau siūlomas vadinamasis algoritminės inžinerinės atsakomybės modelis – analogiškas profesinei atsakomybei statybos ar medicinos srityje, pagal kurį DI kūrėjo atsakomybė būtų grindžiama ne kiekviena konkrečia sprendimo pasekme, bet tuo, kaip sistema buvo suprojektuota, ar buvo atlikta rizikų analizė bei naudotojui suteikta išsami informacija apie logikos ribotumus. Toks modelis analogiškas atsakomybei medicinoje ar statyboje, kai gaminio autorius atsako už konstrukcijos integralumą ir naudojimo saugumą, bet ne už kiekvieną galutinio naudotojo veiksmą. Šią sampratą grindžia ir Čerka, Grigienė bei Sirbikytė (2015), pažymėdami, kad „atsakomybė turi būti diferencijuojama tarp skirtingų subjektų – DI kūrėjo, kuris sukuria algoritminę struktūrą, ir darbdavio, kuris ją taiko darbo santykiuose“ (p. 378). Autoriai siūlo taikyti funkcinį modelį, kuriame kūrėjui tenka griežtoji atsakomybė už logikos konstrukciją, o naudotojui – už konkretaus taikymo sprendimus.

Europos Sąjungos dirbtinio intelekto akte, Reglamente (ES) 2024/1689 – įtvirtinta, kad aukštos rizikos DI sistemų naudotojai, įskaitant darbdavius, privalo užtikrinti, kad sistemos būtų naudojamos pagal paskirtį, būtų taikoma žmonių priežiūra ir įgyvendinta aiški atskaitomybės tvarka (29 str.). Ši nuostata grindžiama principu, jog DI negali pakeisti žmogaus atsakomybės, kai sprendimai daro reikšmingą poveikį asmeniui. Papildomai, Bendrasis duomenų apsaugos reglamentas (Reglamentas (ES) 2016/679) 22 straipsnyje įtvirtina darbuotojo teisę nesutikti su vien automatizuotu sprendimu, jei jis gali turėti reikšmingą arba teisinių pasekmių. Tokiu būdu nacionalinės teisės įpareigojimai darbdaviui turi būti aiškiai derinami su bendraisiais skaidrumo, informavimo ir žmogaus dalyvavimo principais.

Tarptautinėje praktikoje šie principai įgyja konkretų teisinį turinį. Pavyzdžiui, JAV byloje *Mobley v. Workday, Inc.* (Case No. 3:23-cv-00770, N.D. Cal. 2023) buvo sprendžiama, ar algoritminė darbo atrankos sistema, kurią naudojo darbdavys, pažeidžia federalines diskriminacijos normas. Teismas konstatavo, kad net jei sprendimą inicijuoja algoritmas, atsakomybė išlieka darbdaviui pagal Title VII of the Civil Rights Act of 1964. Šis pavyzdys sustiprina nuostatą, jog dirbtinis intelektas nėra savarankiškas atsakomybės subjektas – atsakomybė tenka tam, kas jį taiko praktikoje.

Vienu reikšmingiausių teismų precedentų dėl algoritminės atsakomybės Europoje tapo Nyderlandų byla dėl SyRI (System Risk Indication) sistemos. Ši sistema, sukurta Socialinių reikalų ir užimtumo ministerijos, buvo skirta nustatyti socialinių išmokų sukčiavimą bei kitus pažeidimus pasitelkiant didžiųjų duomenų analizę. 2020 m. Hagos apygardos teismas (Rechtbank Den Haag) sprendime *NJCM c.s. v. Staat der Nederlanden* konstatavo, kad SyRI sistema pažeidė Europos žmogaus teisių konvencijos (EŽTK) 8 straipsnį. Pažeidimas kilo dėl nepakankamos žmogaus priežiūros, informacijos trūkumo apie algoritminę logiką, neproporcingo duomenų kaupimo masto ir neskaidraus veikimo principo. Nors kovos su sukčiavimu tikslas buvo laikomas teisėtu, priemonė – SyRI – nebuvo laikoma proporcinga.

Kaip pažymi Vedder ir Naudts (2020), SyRI veikia kaip „juodoji dėžė“, kurioje sprendimų priėmimo logika lieka neatskleista nei sistemą įgyvendinančioms institucijoms, nei asmeniui, kuriam šis sprendimas taikomas. Tokia algoritminė nepermatomumo struktūra eliminuoja atskaitomybę, o tai reiškia, kad neįmanoma efektyviai apsiginti nuo galimų sprendimo klaidų ar diskriminacinio poveikio (Vedder ir Naudts, 2020, p. 134–136). SyRI byla tapo vienu pirmųjų pavyzdžių Europos teisėje, kai teismas aiškiai pripažino, jog sisteminis, bet nepermatomas ir žmogaus nekontroliuojamas algoritmas pažeidžia pagrindines teises – šiuo atveju, teisę į privatų gyvenimą ir apsaugą nuo savavališko duomenų naudojimo.

Net ir tuo atveju, kai dirbtinio intelekto sistema veikia su socialiai pateisinamu tikslu, tačiau jos veikimas yra neskaidrus, neperžiūrimas arba sukelia reikšmingų pasekmių darbuotojui, galutinė atsakomybė tenka žmogui, darbdaviui. Tokia pozicija grindžiama vadinamuoju atsakingo naudotojo principu, kuris šiuo metu formuojasi tiek nacionalinėje, tiek tarptautinėje teisės doktrinoje. Šis principas reiškia, kad atsakingas yra ne pats DI algoritmas, bet subjektas, kuris jį taiko, leidžia jam veikti ar integruoja į teisinį santykį. Tarptautinė darbo organizacija pažymi, kad „atsakomybė už automatizuotas sistemas turi būti taikoma proporcingai sprendimo galios lygiui, o ne tik formaliai subjekto pozicijai“ (ILO, 2021, p. 12). Rimkutė (2022) papildomai pastebi, kad tokiose situacijose formuojasi „objektyvios kaltės priskyrimo sumaištis“, kai darbuotojas tampa atsakomybės objektu, nors neturėjo galimybės suprasti sprendimo ar jam daryti įtakos (p. 89).

Šio principo įtvirtinimas praktikoje ypač aiškiai atsispindi Vokietijos darbo teisėje. Pagal Betriebsverfassungsgesetz (Vokietijos Darbo tarybų įstatymą), §87 straipsnio 1 dalies 6 punktą, darbo taryba turi bendrą sprendimo teisę, kai įmonėje planuojama diegti technines priemones, skirtas darbuotojų elgesiui arba veiklai stebėti. Tai reiškia, kad jokios DI sistemos, kurios automatizuoja našumo, darbo laiko, judėjimo ar kitus darbo aspektus, negali būti diegiamos be darbo tarybos sutikimo. Šis kolektyvinis kontrolės mechanizmas atlieka esminę darbuotojų teisių užtikrinimo funkciją: jis apsaugo nuo vienpusės darbdavio technologinės galios, įtvirtina pareigą informuoti ir konsultuotis, o taip pat įgalina darbuotojų atstovus veikti kaip institucinę atsvarą algoritminei valdžiai (Todoli-Signes, 2021). Tokia sistema, remiantis ILO rekomendacijomis, užtikrina demokratinį sprendimų priėmimo procesą ir proporcingą atsakomybės paskirstymą DI technologijų kontekste (ILO, 2021, p. 12).

Prancūzijos teisinėje sistemoje darbuotojų atstovų dalyvavimas sprendimų dėl technologinių pokyčių priėmimo laikomas esminiu demokratinio darbo santykių mechanizmo elementu. Pagal *Code du travail* L.2312-8 straipsnį, darbdaviai privalo informuoti ir konsultuotis su ekonomine-socialine taryba (CSE) dėl bet kokios priemonės ar pokyčio, kuris gali paveikti darbo organizavimą, sąlygas ar darbuotojų statusą. Ši pareiga taikoma ir dirbtinio intelekto pagrindu veikiančioms sprendimų priėmimo sistemoms, ypač toms, kurios susijusios su darbo grafikų sudarymu, našumo vertinimu ar automatizuotu sprendimų priėmimu. Kaip nurodoma Prancūzijos Darbo ministerijos gairėse, šis konsultavimosi procesas turi būti vykdomas prieš bet kokią sprendimą ir apimti aiškų paaiškinimą apie technologijos pobūdį, paskirtį bei poveikį darbuotojams. Nesilaikant šios pareigos, CSE gali kreiptis į darbo inspekciją ar teismą, siekdama sprendimo pripažinimo neteisėtu arba jo sustabdymo.

Toks modelis išryškina, kad stiprūs kolektyvinio dalyvavimo mechanizmai atlieka esminį vaidmenį algoritminės kontrolės kontekste: jie užtikrina demokratinį sprendimų priėmimą, skaidrumą ir proporcingą atsakomybės paskirstymą. Jei algoritminiai sprendimai tampa nepermatomi ar taikomi be žmogaus kontrolės, darbuotojų teisės į informavimą ir ginčijimą yra pažeidžiamos.

Todėl nacionaliniu teisėkūros lygmeniu tikslinga būtų aiškiai įtvirtinti DI atsakomybės modelius. Pavyzdžiui, griežtosios atsakomybės modelį DI kūrėjo atžvilgiu. Tai užtikrintų darbdavio galimybę atlyginus žalą darbuotojui, esant pagrindui kreiptis tiesiogiai į DI sistemos kūrėją, tiekėją. Tokiu atveju, darbuotojui taptų paprasčiau ginti pažeistas teises, kai nacionaliniuose teisės aktuose būtų įtvirtinta aiški DI sprendimo priėmimo apskundimo tvarka ir atsakomybės modeliai už DI padarytą žalą darbuotojams. Darbdaviai turėtų būti laikomi atsakingais už DI sistemų diegimą darbo santykiuose ir šių sistemų žalą darbuotojams. Nacionaliniuose aktuose reikalinga įtvirtinti aiškią pareigą konsultuotis DI kontekste prieš

diegiant DI sistemas, darančias poveikį darbo organizavimui, našumo vertinimui ar darbo sąlygoms. Toks reglamentavimas sustiprintų socialinio dialogo mechanizmus, padidintų sprendimų skaidrumą bei sumažintų teisinių ginčų riziką ateityje. Be to, jis užtikrintų, kad technologiniai sprendimai būtų vertinami ne vien pagal efektyvumo kriterijus, bet ir atsižvelgiant į darbuotojų teises ir galimybę dalyvauti sprendimų priėmime. Vokietijos pavyzdys rodo, jog tik tinkamai įtvirtintas kolektyvinis informavimas ir konsultavimasis gali tapti realiu saugikliu prieš DI pagrįstų sprendimų savivalę.

3.2. Civilinės atsakomybės už DI klaidas ir diskriminaciją modeliai

Dirbtinio intelekto integracija į darbo santykių struktūras transformuoja ne tik darbo organizavimo praktiką, bet ir atsakomybės priskyrimo teisinius mechanizmus. DI sistemų sprendimai gali būti grindžiami koreliaciniais algoritmais, kurie nustato ryšius tarp duomenų kintamųjų, tačiau neatskleidžia priežastinio ryšio. Tai reiškia, kad sprendimai priimami autonomiškai arba pusiau autonomiškai, be žmogaus tiesioginio valios išreiškimo, todėl tradicinės civilinės atsakomybės konstrukcijos – paremtos aiškia neteisėto veiksmo, kaltės, žalos ir priežastinio ryšio grandimi – tampa problemiškos. Lietuvos Respublikos civilinio kodekso VI knygos ketvirtasis skirsnis „Civilinė atsakomybė“ (CK 6.246–6.249 str.) reikalauja nustatyti visus šiuos elementus kaip sąlygą atsakomybei kilti. Kaip pažymi Bitinas, Tartilas ir Litvaitienė (2011), būtent šie kriterijai sudaro deliktinės atsakomybės pagrindą, tačiau algoritminių sprendimų aplinkoje jie dažnai išnyksta arba tampa sunkiai identifikuojami. Murauskas pagrįstai atkreipia dėmesį, kad algoritmų taikymo vertinimas teisinėje aplinkoje negali būti atliekamas abstrakčiai – jis turi būti susietas su konkrečiomis sprendimo priėmimo stadijomis. Tik toks klasifikacinis požiūris leidžia identifikuoti rizikas bei taikyti proporcingus teisinius saugiklius. Autorius siūlo naudoti trijų stadijų modelį – faktų vertinimo, taikytinos teisės nustatymo ir sprendimo priėmimo – pagal kurį galima suprasti, kokio pobūdžio algoritmai veikia atitinkamu etapu ir kokių apsaugos priemonių reikia. Teisė, pasak Murausko, tokiu atveju turi peržengti formalų schematiškumą ir įvertinti faktinę žmogaus galimybę kontroliuoti algoritminį sprendimą (Murauskas, 2020, p. 68).

Atsakomybės negalima grįsti vien formaliu subjektų pasiskirstymu pagal teisinį statusą, pavyzdžiui, darbdavio ar techninės sistemos tiekėjo vaidmenimis, nes DI sprendimai dažnai peržengia klasikinės subjektyvios kaltės sampratą. Lietuvos Respublikos civilinis kodeksas numato bendrąją atsakomybę už padarytą žalą, kuri, kaip nustatyta 6.263 straipsnio 1 dalyje, gali kilti ir be kaltės, kai to reikalauja įstatymai arba sutartis. Tai tampa ypač reikšminga, jei DI veikia savarankiškai ir nėra aiškaus žmogaus valios pėdsako. Europos Komisijos 2022 m. Dirbtinio intelekto atsakomybės direktyvos pasiūlyme (COM(2022) 496 final, p. 6) akcentuojama, kad „atsižvelgiant į technologijos autonomiškumą ir veikimo nenusipėjamumą, griežtoji atsakomybė tampa proporcinga priemone pažeistoms teisėms apginti“. Šis požiūris suponuoja, kad DI naudotojai ir kūrėjai negali būti atleisti nuo atsakomybės vien dėl to, kad sprendimą suformavo ar padėjo formuoti algoritmas – būtinas objektyvus, veiklos rizika grindžiamas požiūris, žmogaus įsikišimas. Algoritminiai sprendimai veikia pagal statistinius modelius, kurie generuoja sisteminės, o ne atsitiktinės klaidas, todėl jų pasekmės dažnai negali

būti siejamos su konkretaus subjekto kalte. Net kai pažeidimą lemia techninis DI sistemos sutrikimas, teisės aktuose nėra aiškių normų, kaip paskirstyti atsakomybę tarp technologijų kūrėjo, sistemos naudotojo (darbdavio). Lietuvos civilinis kodeksas numato bendrąją atsakomybę už padarytą žalą (CK 6.263 str.), tačiau DI kontekste kyla klausimas, kas yra „atsakingas subjektas“, kai sprendimus lemia iš anksto apmokytas, vėliau autonomiškai veikiantis algoritmas. Kaip numatyta CK 6.263 straipsnio 1 dalyje, atsakomybė už žalą gali būti taikoma ir be kaltės, jei tai numato įstatymas arba sutartis. DI kontekste ši nuostata tampa ypač reikšminga, jei sprendimai būtų priimami be žmogaus tiesioginio įsikišimo. Scholz (2020, p. 251) pagrįstai teigia, kad „atsakomybės paskirstymas tarp žmogaus ir DI tampa nebe teisės klausimu, o moralinio ir socialinio pagrįstumo dilema“. Ši dilema dar labiau komplikuojama Europos Parlamento ir Tarybos reglamento (ES) 2024/1689 (Dirbtinio intelekto akto) 29 straipsnio nuostatomis, kurios naudotojams, tarp jų ir darbdaviams, numato pareigą užtikrinti žmogaus priežiūrą, paaiškinamumą bei pagrindinių teisių laikymąsi diegiant DI sistemas. Todėl civilinės atsakomybės diskusijoje būtina taikyti funkcinį atsakomybės modelį, kuriame įvertinamas ne tik formaliai priskirtas statusas, bet ir faktinis dalyvavimas priimant sprendimus bei gebėjimas juos kontroliuoti. Rizikos paskirstymo principas tokiu atveju reiškia, kad daugiau galios turintys subjektai turi prisiimti didesnę atsakomybę už sprendimo pasekmes.

Darbuotojų atsakomybė algoritminės kontrolės sąlygomis turėtų būti vertinama atsižvelgiant į jų realią įtaką sprendimo priėmimui. Kaip pažymi Grumulaitis (2023, p. 4), „darbuotojas, kuris veikia tik kaip sistemos vykdytojas, negali būti laikomas atsakingu už pasekmes, kylančias iš algoritminio sprendimo, kurio veikimo principas jam nežinomas ar nepaaiškintas.“ Ši pozicija atitinka Tarptautinės darbo organizacijos (ILO) rekomendacijas, kuriose pabrėžiama, kad atsakomybė už automatizuotas sistemas turi būti priskiriama proporcingai sprendimų galios lygiui, o ne vien pagal formalų subjekto statusą (ILO, 2021, p. 12). Toks principas dera su nacionalinėje teisėje taikomu atsakomybės personalizavimo kriterijumi, asmuo negali būti laikomas atsakingu be realaus priežastinio ryšio tarp jo veiksmų ir padarinių. Kaip rodo Lietuvos socialinių tyrimų centro duomenys, algoritminė kontrolė darbo aplinkoje darbuotojams dažnai yra ne tik nepermatoma, bet ir sukelia neužtikrintumo, asmeninio poveikio stokos bei atsakomybės baimės jausmą, ypač kai sprendimai priimami automatiškai be galimybės juos suprasti ar paveikti (LSTC, 2023, p. 82–83).

Darbo teisė, nors ir laikoma savarankiška teisės šaka, istoriškai ir sistemiškai yra glaudžiai susijusi su civiline teise, ypač sutarčių ir deliktų institutais. Kaip pažymi Davulis (2008, p. 28), darbo teisė išsivystė iš civilinės teisės pagrindo, tačiau dėl augančio viešojo intereso, socialinės politikos prioritetų ir darbuotojo kaip silpnesnės darbo santykių šalies apsaugos poreikio ji išgrynino savarankiškumą. Vis dėlto ši autonomija nesunaikina civilinės

teisės instrumentų aktualumo, ypač kai darbo teisės normos neapima technologinių naujovių sukeltų situacijų, pavyzdžiui, kai darbo aplinkoje veikia dirbtinio intelekto sprendimų priėmimo sistemos, veikiančios be žmogaus įsikišimo arba nepriklausomai nuo darbdavio intencijų. Tokiais atvejais, vadovaujantis subsidiarumo principu, įtvirtintu Lietuvos Respublikos darbo kodekso 4 straipsniu, taikomos civilinės teisės normos. Ypač reikšminga tampa deliktinė atsakomybė, kuri veikia ne tik kaip kompensacinė priemonė nukentėjusiajam, bet ir kaip prevencinis mechanizmas, skatinantis darbdavius bei DI kūrėjus diegti etiškai ir teisiškai atsakingas technologijas. Tokia sąveika tarp darbo ir civilinės teisės įgalina sistemiškai spręsti sudėtingas situacijas, kai dirbtinis intelektas sukuria rizikas, kurių darbo teisės institucinė struktūra tradiciškai nėra pajėgi aprėpti.

Vienas reikšmingiausių precedentų Europoje buvo Hagos apygardos teismo sprendimas byloje *NJCM c.s. v. Staat der Nederlanden dėl SyRI* (System Risk Indication) sistemos. Teismas 2020 m. konstatavo, kad valstybės įdiegta socialinės rizikos vertinimo sistema, veikusi be aiškios sprendimo logikos ir neproporcingai kaupusi duomenis, pažeidė Europos žmogaus teisių konvencijos 8 straipsnį – teisę į privatų gyvenimą. Tai buvo vienas pirmųjų atvejų, kai algoritminis netransparentumas teisiškai įvertintas kaip žmogaus teisių pažeidimas. Nors ši byla buvo susijusi su viešojo administravimo sritimi, jos išvados turi reikšmingų pasekmių ir privačiame sektoriuje, ypač darbo teisėje. Jei darbdavys naudoja algoritminę sistemą, kuri lemia teisiškai reikšmingus sprendimus (pvz., dėl atleidimo), tačiau jos veikimas darbuotojui yra nepermatomas, tai gali kelti ne tik duomenų apsaugos ar nediskriminavimo problemas, bet ir civilinės atsakomybės. Tokiu atveju darbdavys galėtų atsakyti ne tik pagal DK ribas, bet toks sprendimas galėtų būti vertinamas ir kitų teisės aktų kontekste.

Kaip pažymi Juškevičiūtė-Vilienė (2020), „dirbtinio intelekto sprendimai gali sukelti konstitucinės teisės į ginčą ir teisingumą pažeidimą, jei asmuo neturi prieigos nei prie sprendimo pagrindimo, nei jo logikos“ (p. 101). Tokia jurisprudencija ir doktrina rodo, kad net ir formaliai automatizuoti sprendimai, inicijuoti DI sistemų, turi būti vertinami pagal žmogaus teisių, skaidrumo ir atskaitomybės principus. Tai savo ruožtu suponuoja būtinybę taikyti civilinės atsakomybės mechanizmus, net jei žmogaus valia sprendimo priėmimo procese buvo minimali ar neegzistavo – nes sprendimo pasekmės vis tiek paveikia darbuotojo teisinį statusą, teisę į veiksmingą gynybą.

Lietuvos Respublikos darbo kodekso 4 straipsnio 2 dalyje įtvirtintas subsidiarumo principas, pagal kurį, jei tam tikri darbo santykių aspektai nėra tiesiogiai reglamentuoti darbo teisėje, taikomos civilinės teisės normos tiek, kiek jos neprieštarauja darbo teisės tikslams ir principams. Tokia situacija gali kilti, kai darbdavys naudoja automatizuotas sprendimų priėmimo sistemas, grindžiamas dirbtinio intelekto (DI) technologijomis, tačiau darbo teisėje trūksta aiškių

nuostatų dėl šių sistemų paaiškinamumo ar jų sukeltos žalos kompensavimo mechanizmų. Pavyzdžiui, jei darbuotojas patiria neturtinę žalą dėl automatizuoto vertinimo ar sprendimo, kuris pažeidžia jo orumą ar lygias galimybes, o jei tokia situacija nėra tiesiogiai reglamentuota DK, pagrindas ginti teises gali būti randamas Civilinio kodekso normose – ypač CK 6.250 straipsnio 2 dalyje, kuri leidžia priteisti neturtinę žalą, kai pažeidžiamos asmens teisės ir dėl to atsiranda dvasiniai išgyvenimai, pažeminimo ar nesaugumo jausmas. Nors Civilinis kodeksas tiesiogiai nenumato automatizuoto sprendimo kaip žalos šaltinio, tam tikromis aplinkybėmis, kai darbuotojui nesuteikiama galimybė suprasti sprendimo logikos, paaiškinti jo reikšmės ar jį ginčyti, o tai sukelia reikšmingą psichologinį diskomfortą, gali būti pagrindas taikyti deliktinės atsakomybės modelį. Tokiais atvejais, pareiga pagrįsti sprendimą įgauna ne tik procedūrinę, bet ir materialinę reikšmę asmens teisių apsaugos kontekste. Teisinio tikrumo ir orumo principai suponuoja, kad darbuotojas turi teisę į prognozuojamus, paaiškinamus ir žmogaus kontroliuojamus sprendimus, net jei jie priimami taikant technologinius įrankius.

Situacijos, kai darbuotojui padaroma žala dėl automatizuoto sprendimo, priimto be žmogaus įsikišimo ar be pakankamo pagrindimo, kelia klausimą dėl neturtinės žalos atlyginimo pagal civilinės teisės normas. Pavyzdžiui, kai dirbtinio intelekto sistema, naudojama darbuotojų našumui vertinti, suformuluoja klaidingą ar diskriminuojantį sprendimą dėl amžiaus, lyties ar negalios, darbuotojas gali patirti profesinę, reputacinę ar emocinę žalą. Tokiais atvejais, kai sprendimo logika yra nepermatoma (angl. *opaque*), o darbdavys neįvykdo pareigos informuoti ar paaiškinti sprendimą, susiformuoja teisinė padėtis, kurioje galima svarstyti neturtinės žalos atlyginimą pagal Civilinio kodekso 6.250 straipsnį. Neturtinė žala tokiu atveju kyla iš žmogaus lūkesčių ir teisės į paaiškinimą pažeidimo. Pareiga informuoti, pagrįsti sprendimą ir užtikrinti žmogaus įsikišimą aiškiai išplaukia tiek iš Bendrojo duomenų apsaugos reglamento (Reglamentas (ES) 2016/679, 13 ir 22 str.), tiek iš Dirbtinio intelekto akto (Reglamentas (ES) 2024/1689, 29 str.). Šie teisės aktai nustato, kad darbuotojo teisė žinoti sprendimo pagrindą turi būti ne deklaratyvi, bet realiai įgyvendinama.

Kai darbuotojas priima sprendimus naudodamasis darbdavio suteiktais dirbtinio intelekto pagrindu veikiančiais įrankiais, kyla sudėtingas atsakomybės paskirstymo klausimas ypač tais atvejais, kai dėl tokio sprendimo žala padaroma trečiajam asmeniui, pavyzdžiui, klientui. Pagal Lietuvos Respublikos civilinio kodekso 6.264 straipsnį, už darbuotojo, einančio savo pareigas, padarytą žalą atsako darbdavys kaip juridinis asmuo, nepriklausomai nuo to, ar darbuotojas veikė pagal tiesioginį nurodymą, ar naudodamasis jam suteiktais instrumentais. Tokiu atveju klientas turi teisę reikšti ieškinį darbdaviui, kuris atsako kaip pavaldinio veiksmų subjektas (CK 6.264 str. 1 d.). Tuo tarpu darbuotojo atsakomybė darbdaviui, kai šis dėl pavaldinio veiksmų patiria žalą (pavyzdžiui, dėl netinkamo DI taikymo ar savavališko sprendimo priėmimo), reglamentuojama

Darbo kodekso 154 ir 155 straipsniais. Pagal DK 154 str., darbuotojas atsako už žalą tik tuo atveju, kai ji padaryta dėl jo kaltės, o DK 155 str. nustato, kad darbuotojas atlygina žalą tik tuo atveju, kai ji padaryta netinkamai vykdant pareigas, kai jis nesilaikė nustatytos darbo tvarkos ar viršijo savo kompetencijos ribas. Taigi, jeigu darbuotojas sąžiningai taikė darbdavio suteiktą DI sistemą, atsakomybė tenka darbdaviui. Priešingai, jeigu darbuotojas savo iniciatyva pasitelkė DI sprendimą be darbdavio žinios ar leidimo, gali būti sprendžiama dėl jo atsakomybės pagal DK 154–155 str., priklausomai nuo kaltės laipsnio ir padarytos žalos masto.

Situacijose, kurios nėra detalios reglamentuotos Darbo kodekse, pavyzdžiui, kai DI sprendimai daro poveikį ne tik darbdaviui, bet ir išorės subjektams, aktualus tampa DK 4 straipsnio 2 dalyje įtvirtintas subsidiarumo principas, leidžiantis taikyti civilinės teisės normas darbo teisės spragoms užpildyti. Tokiu būdu CK 6.264 str. gali būti taikomas tiek darbuotojo atsakomybei trečiajam asmeniui, tiek darbdavio priklausomai nuo faktinių aplinkybių ir kompetencijos pasiskirstymo.

Scholz (2020, p. 251) pažymi, kad algoritmų naudojimas neišlaisvina darbdavio nuo pareigos užtikrinti ne tik formalią teisės aktų atitikimą, bet ir darbo santykių moralinį bei socialinį pagrįstumą. Ši mintis ypač svarbi taikant Europos Sąjungos Dirbtinio intelekto akto 29 straipsnį, kuris įpareigoja naudotojus, įskaitant darbdavius, užtikrinti žmogaus priežiūrą, sprendimų skaidrumą ir pagrindinių teisių laikymąsi diegiant aukštos rizikos DI sistemas. Tokia prievolė transformuoja darbdavio vaidmenį iš pasyvaus naudotojo į aktyvų atsakingą subjektą, turintį užtikrinti, kad technologijos būtų diegiamos etiškai ir teisiškai atsakingai.

Todėl darbdavio atsakomybė DI kontekste neturi būti vertinama siaurai – ji apima ne tik sprendimo priėmimą, bet ir veikimo sąlygų sukūrimą, pasekmių valdymą bei teisinių garantijų įgyvendinimą. Ši atsakomybės struktūra pagrįsta ne tik subjektyvia kalte, bet ir objektyvia pareiga užtikrinti teisėtą technologijos veikimą darbo santykių kontekste.

Nors darbo teisė pirmiausia orientuota į viešojo intereso apsaugą bei darbuotojo teisių užtikrinimą per institucinę sistemą, civilinė teisė remiasi individualių teisių ir interesų balansu bei kompensacijos principu. Algoritminiai sprendimai darbo aplinkoje dažnai formaliai priskiriami darbdaviui, tačiau jų veikimas grindžiamas technologiniais modeliais, kurių sprendimų logika gali būti neaiški, kintanti arba visiškai neprieinama. Tokiais atvejais, kai darbuotojas patiria žalą ir tokios situacijos dėl patirtos žalos nėra DK reguliavimo dalykas, darbuotojas negali įvertinti ar užginčyti sprendimo pagrįstumo, atsakomybė tampa civilinės teisės reguliavimo objektu. Civilinė atsakomybė tokiu būdu tampa būtinu kompensaciniu mechanizmu, kai institucinės darbo teisės garantijos pasirodo esančios nepakankamos.

DI sistemos dažnai veikia remdamasis sudėtingais duomenų apdorojimo modeliais, kurie leidžia sistemai „numatyti“ tam tikrą rezultatą pagal ankstesnius panašius atvejus, bet neparodo

aiškos priežasties, kodėl sprendimas buvo priimtas būtent taip. Kitaip tariant, sistema pateikia rezultatą pagal tam tikrų požymių sutapimus, bet ne remdamasi konkrečia ir teisiškai įrodomai nustatyta sprendimo logika. Dėl to tampa sunku nustatyti, ar buvo padarytas neteisėtas veiksmas, kas už jį atsakingas ir kaip vertinti kaltės buvimą. Kaip pažymi Čerka, Grigienė ir Sirbikytė (2015, p. 378), „klasikinės deliktinės atsakomybės schemos žlunga ten, kur neįmanoma atriboti techninės algoritminės logikos nuo žmogaus valios, nes ši logika vystosi *post factum*“. Tokį požiūrį palaiko ir Europos Komisijos parengtas Dirbtinio intelekto atsakomybės direktyvos pasiūlymas, kuriame pabrėžiama, kad „atsižvelgiant į technologijos autonomiškumą ir veikimo nenuspėjamumą, griežtoji atsakomybė tampa proporcinga priemone pažeistoms teisėms apginti“ (Europos Komisija, 2022, p. 6).

Nors praktikoje darbdavys, kaip DI naudotojas, dažniausiai laikomas atsakingu už sprendimus darbo santykiuose, ši pozicija neišsprendžia visų civilinės atsakomybės problemų. Nors sprendimą formaliai priima darbdavys, faktinė kontrolė dažnai pasiskirsto tarp technologinės sistemos elementų, ypač tais atvejais, kai naudojami adaptyvūs arba savimokos modeliai. Tokiose situacijose klasikinė kaltės samprata tampa problemiška. Be to, kaip nurodo Tarptautinė darbo organizacija, atsakomybė už automatizuotas sistemas turi būti proporcinga sprendimų kontrolės lygiui ir galimybei paaiškinti jų eigą (ILO, 2021, p. 25). Tokia pozicija atspindi šiuolaikinių teisės doktrinų tendenciją pereiti prie funkcinio atsakomybės modelio, kuriame subjektas atsako ne pagal formalią poziciją, o pagal realų dalyvavimo pobūdį. ES Dirbtinio intelekto aktas (Reglamentas (ES) 2024/1689, 29–30 str.) taip pat įpareigoja naudotojus užtikrinti žmogaus priežiūrą, skaidrumą ir sprendimo paaiškinamumą. Todėl deliktinės atsakomybės analizėje būtina peržengti formalių kriterijų ribas ir įvertinti faktinę galios struktūrą, ypač kai sprendimo padariniai turi reikšmingą poveikį darbuotojo teisėms.

Atsižvelgiant į tai, atsakomybės modelis DI kontekste turėtų būti grįstas ne vien klasikine deliktinės atsakomybės sąlygomis, bet ir sisteminiu požiūriu, įvertinant technologinės architektūros įtaką sprendimų priėmimui. Klasikinis deliktas remiasi subjektyvia žmogaus valia, tačiau algoritminių sistemų sprendimai dažnai kuriami ir įgyvendinami be aiškos žmogaus intervencijos. Tokiais atvejais tampa būtina diferencijuoti atsakomybę pagal dalyvavimo pobūdį: DI kūrėjui – griežtoji atsakomybė už algoritminės logikos struktūrą, naudotojui – atsakomybė už sprendimo taikymo kontekstą. Kaip pažymi Wachter, Mittelstadt ir Floridi (2017), teisė į prieigą prie sprendimo logikos yra neatsiejama nuo asmens galimybės efektyviai apsiginti, o jos ignoravimas „atima materialų turinį iš teisės ginčyti sprendimą“ (p. 497). Todėl funkcinės atsakomybės modelis suteikia pagrindą sprendimų vertinimui ne pagal formalią hierarchiją, o pagal faktinę įtaką rezultatui. Europos Komisijos DI atsakomybės direktyvos projektas taip pat siūlo pereiti prie tokio modelio, įtvirtinant galimybę darbuotojui

reikalauti atsakomybės tiek iš naudotojo, tiek iš technologijos tiekėjo, priklausomai nuo jų įnašo į sprendimo turinį (Europos Komisija, 2022, p. 8).

Be to, Europos Sąjungos Dirbtinio intelekto reglamentas (Reglamentas (ES) 2024/1689) numato pareigą DI naudotojams užtikrinti žmogaus priežiūrą, sprendimų paaiškinamumą ir atitiktį pagrindinėms teisėms (str. 29–30). Tai patvirtina, kad atsakomybė negali būti vienpusė, civilinės teisės normos būtinos tam, kad būtų galima reikšti pretenzijas ir DI kūrėjui ar platintojui, jei, pavyzdžiui, algoritmas jau iš anksto įdiegtas su diskriminuojančiomis prielaidomis.

Wachter, Mittelstadt ir Floridi (2017) nurodo, kad vien teisė nesutikti su automatizuotu sprendimu pagal BDAR 22 straipsnį nepakankama, jei duomenų subjektui nesuteikiama reali galimybė suprasti sprendimo logiką ar gauti veiksmingą žalos atlyginimą. Jie argumentuoja, kad esant tokiai situacijai, būtina stiprinti civilinės atsakomybės mechanizmus, kad būtų užtikrintas materialus teisės į paaiškinimą ir gynybą įgyvendinimas. Todėl, nors darbdavys ir išlieka pagrindiniu atsakomybės centru darbo teisės prasme, civilinės teisės normos gali būti būtinos siekiant užtikrinti proporcingą atsakomybės paskirstymą, įtraukiant ne tik naudotojus, bet ir technologijų kūrėjus bei tiekėjus. Be tokio reguliavimo, darbuotojas lieka be veiksmingų priemonių ginčyti ar kompensuoti algoritminės žalos padarinius.

Pavyzdžiui, byloje *Mobley v. Workday Inc.* (JAV, 2023), nagrinėtas atvejis, kai automatizuota atrankos sistema galimai diskriminavo vyresnio amžiaus kandidatus. Teismas atkreipė dėmesį, kad naudotojas negali atsiriboti nuo atsakomybės vien dėl to, jog sprendimą suformulavo išorinė technologinė sistema, o ne jis pats – todėl jam tenka pareiga užtikrinti atitiktį diskriminacijos draudimo normoms (*Mobley v. Workday Inc.*, 2023 WL 5094191). Praktikoje vis dažniau fiksuojami atvejai, kai automatizuoti sprendimai sukelia diskriminacines pasekmes, ypač kai jie grindžiami neišaiškinama logika arba koreliaciniais duomenimis, kurie nesuvokiami darbuotojui. Tokios situacijos patenka į Bendrojo duomenų apsaugos reglamento (BDAR) 22 straipsnio ir Dirbtinio intelekto akto 29 straipsnio taikymo sritį, įtvirtinančią naudotojo pareigą užtikrinti paaiškinamumą ir apsaugą nuo šališkumo.

Civilinės atsakomybės institutas, kuris tiesiogiai susijęs su darbuotojo ar darbdavio patirta žala dirbtinio intelekto kontekste reikalauja esminės paradigmos peržiūros – nuo subjektyvios kaltės nustatymo prie funkcinio atsakomybės paskirstymo, atsižvelgiant į technologinę architektūrą ir sprendimo įtaką darbuotojo teisėms. Nors darbo teisėje pagrindinė atsakomybė už darbuotojo teisių apsaugą tenka darbdaviui, vien šios atsakomybės nepakanka. Siekiant teisingo ir veiksmingo atsakomybės paskirstymo dirbtinio intelekto aplinkoje, būtina įtraukti ir kitus grandinės dalyvius – technologijų kūrėjus bei tiekėjus. Civilinės teisės normos čia atlieka papildomą, bet esminį vaidmenį, nes tik jos leidžia sukurti išplėstinį, kelių šaltinių

pareigomis grindžiamą atsakomybės modelį. Jei toks modelis nėra įtvirtintas, darbuotojas atsiduria teisinio vakuumo situacijoje – be realių galimybių apskųsti arba atlyginti algoritminio sprendimo sukeltą žalą. Nacionalinės teisės reguliavimo galimybės dažnai pasirodo nepakankamos automatizuotų sprendimų tarpvalstybinio poveikio kontekste. Tokiu atveju įgyja svarbą tarptautinės teisės priemonės, ypač konvencijos, kurios gali užtikrinti nuoseklų atsakomybės reguliavimą tarp skirtingų jurisdikcijų. Šiame kontekste itin reikšmingas tampa Vilniaus konvencijos vaidmuo.

3.3. Vilniaus konvencijos poveikis atsakomybės teisiniam reguliavimui

Naujausia Europos Audito Rūmų 2024 m. ataskaita apie dirbtinį intelektą parodė, kad Europos Sąjunga yra išsikėlusį ambicingus tikslus šioje srityje, tačiau jų įgyvendinimui trūksta nuoseklaus finansavimo, stipresnio koordinavimo ir aiškesnio teisinio pagrindo (Europos Taryba, 2024). Vienas svarbiausių iššūkių, kad reguliavimas nebespėja su technologijų plėtra, o tai ypač matoma darbo teisės kontekste. Vis dažniau darbo santykiuose pasirodo algoritmai, kurie padeda priimti sprendimus dėl darbuotojų. Pavyzdžiui, skirsto darbo užduotis, vertina našumą ar net prisideda prie sprendimų dėl atleidimo. Tačiau daugeliu atvejų nėra aišku, kas atsakingas už šiuos sprendimus, kokiomis taisyklėmis jie remiasi, ir ar darbuotojai turi galimybę ginčyti logiką (ibid., p. 20). Kaip jau buvo analizuota ankstesniame poskyryje, nacionaliniai atsakomybės modeliai, ypač grindžiami subjektyvios kaltės elementais, tampa riboti algoritmų sprendimų kontekste. Europos Tarybos Pagrindų konvencijoje dėl dirbtinio intelekto, žmogaus teisių, demokratijos ir teisinės valstybės, pasirašytoje 2024 m. rugsėjo 5 d. Vilniuje, įtvirtinti principai ir teisiniai reikalavimai gali sisteminiu būdu išplėsti atsakomybės ribas, ypač darbo teisės aplinkoje. Ši konvencija – pirmasis tarptautinis teisiškai privalomas susitarimas, skirtas užtikrinti, kad dirbtinio intelekto sistemos būtų plėtojamoms žmogaus teisių ir demokratiškos vertybių pagrindu. Konvencijoje pabrėžiamas skaidrumo, paaiškinamumo ir atskaitomybės principų svarbos taikymas visame DI sistemų gyvavimo cikle, kas ypač aktualu darbo santykių kontekste, kai sprendimai dėl darbuotojų priimami remiantis algoritmais, dažnai be aiškos žmogaus intervencijos (Europos Taryba, 2024).

Taigi, toks neapibrėžtumas kelia riziką, kad darbuotojų teisės gali būti pažeistos ne dėl piktavališkumo, o tiesiog dėl to, kad teisė kol kas nesukūrė tinkamų mechanizmų tokioms situacijoms reguliuoti. Todėl vis dažniau kalbama apie tai, kad darbo teisė turi ne tik reaguoti į jau įvykusius pažeidimus ar ginčus, bet ir iš anksto numatyti rizikas, kurias sukelia naujos technologijos. Kitaip tariant, reikia ne tik saugoti tai, kas jau aišku, bet ir išplėsti teisės galimybes atpažinti naujas situacijas – ypač tokias, kurias lemia sprendimus priimančios sistemos, o ne žmonės.

2024 m. rugsėjo 5 d. pasirašyta Europos Tarybos Pagrindų konvencija dėl dirbtinio intelekto, žmogaus teisių, demokratijos ir teisinės valstybės (CETS Nr. 225), plačiai vadinama Vilniaus konvencija, žymi naują etapą tarptautiniame DI reguliavime. Tai pirmasis tarptautinis teisiškai privalomas dokumentas, kuriuo siekiama užtikrinti, kad dirbtinio intelekto sistemų veikla visame jų gyvavimo cikle būtų visiškai suderinta su žmogaus teisėmis, demokratija ir teisinės valstybės principais. Konvencija apima ne tik duomenų apsaugą, bet ir atsakomybės,

žalos kompensavimo, skaidrumo bei paaiškinamumo klausimus, todėl jos poveikis nacionaliniams atsakomybės režimams, ypač civilinėje ir darbo teisėje, tampa vis aiškesnis.

Vienas iš esminių šios konvencijos principų, atskaitomybė (angl. *accountability*) – reikalauja, kad bet kuris dirbtinio intelekto sistemos naudotojas ar kūrėjas privalėtų įrodyti, jog jų sprendimų poveikis atitinka pagrindinių teisių apsaugos standartus. Konvencijos 14 straipsnyje įtvirtinta, kad valstybės turi sukurti aiškias, veiksmingas ir prieinamas priemones ginti teises, jei jos pažeidžiamos naudojant DI, įskaitant teisę į sprendimo paaiškinimą bei teisę į veiksmingą teisinę gynybą. Ši nuostata tiesiogiai susijusi su darbo santykiais, kuriuose vis dažniau sprendimus įtakoja algoritmai, tačiau darbuotojas neturi realios galimybės suprasti, kokių pagrindų buvo atliktas vertinimas ar atsirado neigiamas pasekmių grandinės rezultatas.

Vilniaus konvencijoje taip pat pabrėžiama, kad valstybės narės privalo užtikrinti, jog egzistuočių „civilinės teisinės atsakomybės režimai, leidžiantys nukentėjusiam asmeniui gauti žalos atlyginimą, kai žala padaroma dėl dirbtinio intelekto sistemų naudojimo ar jų poveikio“, net ir tais atvejais, kai sprendimas nebuvo tiesiogiai priimtas žmogaus (Europos Taryba, 2024, 6 str.). Ši nuostata atspindi aiškią norminę kryptį, nebe konkretus žmogaus veiksmas, o pati sisteminė pasekmė, kylanti iš DI sistemos veikimo, gali tapti atsakomybės pagrindu. Tai ženkliai keičia tradicinį deliktinės atsakomybės modelį, kuris ilgą laiką buvo grįstas neteisėtu asmens elgesiu, kaltės prezumpcija ir subjektyvumu.

Tokiu būdu konvencija konceptualiai artėja prie objektyvios atsakomybės modelių, kur atsakomybė kyla dėl pavojingo ar socialiai reikšmingo veikimo nepriklausomai nuo kaltės. Šis poslinkis yra būtinas atsižvelgiant į DI sistemoms būdingus bruožus, autonomiškumą, sprendimo grandinės nepermatomumą (angl. *opacity*) ir rezultatų nuspėjamumą. Taip konvencija siūlo transformacinį požiūrį į atsakomybę, kuris gali būti esminis reguliuojant algoritminio sprendimų priėmimo taikymą darbo santykiuose.

Pažymėtina, kad Vilniaus konvencija reikalauja, jog visos sistemos, turinčios reikšmingą poveikį žmogaus teisėms, būtų vertinamos rizikos požiūriu, įskaitant poveikį nediskriminavimo, gyvenimo ir privatumo teisėms. Kadangi dauguma darbo teisės atvejų patenka į šią kategoriją, šis aspektas įpareigoja valstybę sukurti horizontalųjį atsakomybės paskirstymą, apimančią tiek darbdavį, tiek trečiąsias šalis – technologijų kūrėjus ir tiekėjus. Šiame kontekste atsakomybė tampa kolektyvine, paskirstyta ir daugiasluoksni, tam reikia tiek civilinių, tiek darbo teisės instrumentų, nes nė viena šaka atskirai negali aprėpti visų DI sąlygotų santykių struktūrų.

Todėl Vilniaus konvencija, nors formaliai nėra darbo teisės ar civilinės atsakomybės teisės aktas, iš esmės veikia kaip konstitucinis principų šaltinis, praplečiantis atsakomybės modelių sisteminį horizontą: nuo individualaus delikto – iki sisteminės pareigos valdyti rizikas ir užtikrinti kompensacijos galimybes net DI autonominių sprendimų atvejais.

Vienas esminių praktinių klausimų, kylančių Vilniaus konvencijos kontekste, yra tai, kaip valstybės narės turės reformuoti nacionalinius atsakomybės modelius, kad šie atitiktų konvencijoje įtvirtintą priežastinės atsakomybės (angl. *causal liability*) paradigmą. Europos Tarybos Pagrindų konvencijoje dėl dirbtinio intelekto, žmogaus teisių, demokratijos ir teisinės valstybės (CETS Nr. 225) numatyta, kad valstybės narės privalo įtvirtinti „civilinės teisinės atsakomybės režimus“, leidžiančius asmenims reikalauti žalos atlyginimo už DI sistemų poveikį, net jei sprendimą priėmė ne žmogus (Europos Taryba, 2024, 6 str.).

Kaip nurodo Bitinas, „civilinei atsakomybei taikyti būtinas ne tik žalos faktas ir neteisėtas veikimas, bet ir subjektyvusis elementas – kaltė“ (Bitinas et al., 2011, p. 252). Tokia struktūra veikia klasikinėse situacijose, kai žala kyla iš aiškiai identifikuojamo asmens veikimo ar neveikimo. Tačiau dirbtinio intelekto kontekste tradiciniai civilinės atsakomybės kriterijai tampa riboti. DI sistemoms būdingi požymiai, autonomiškumas, sprendimų nepermatomumas (angl. *black-box opacity*) ir algoritminis išankstinis šališkumas, reikšmingai apsunkina subjektyviųjų atsakomybės sąlygų, ypač kaltės, taikymą. Kaltės nustatymas tokiose sistemose neretai tampa neįmanomas ne dėl proceso kokybės stokos, o dėl pačios technologinės struktūros, kai sprendimas nėra aiškiai priskirtinas konkrečiam asmeniui. Tai ne tik komplikuoja nukentėjusio asmens teisę į veiksmingą gynybą, bet ir kelia grėsmę pačios atsakomybės instituto efektyvumui.

Šiame kontekste alternatyviais instrumentais tampa priežastinė arba net objektyvioji atsakomybė, kuri leidžia reaguoti į DI poveikį nepriklausomai nuo subjektyvios kaltės. Vilniaus konvencijoje – t. y. Europos Tarybos Pagrindų konvencijoje dėl dirbtinio intelekto, žmogaus teisių, demokratijos ir teisinės valstybės – įtvirtinta nuostata, kad valstybės narės turi „užtikrinti veiksmingas atsakomybės formas net ir nesant tiesioginio žmogaus sprendimo“, o tai iš esmės atitinka objektyviosios atsakomybės modelį (Europos Taryba 2024, 6 straipsnis). Tai žymi doktrininius pokyčius, nes objektyvioji atsakomybė iki šiol Lietuvos teisėje buvo taikoma tik išimtinėse situacijose – pavyzdžiui, pagal Civilinio kodekso 6.270 straipsnį, susijusį su padidinto pavojaus šaltiniais (transporto priemonės, mechanizmai, gamyklos ir pan.).

Kaip pažymima 6 straipsnyje, kiekviena valstybė turi imtis būtinų priemonių, kad būtų užtikrintos veiksmingos teisių gynimo priemonės, įskaitant tinkamą žalos atlyginimą ir kompensavimą, kai dirbtinio intelekto sistemų naudojimas daro neigiamą poveikį žmogaus teisėms (Europos Taryba 2024, 6 straipsnis). Ši formuluotė rodo, kad pareiga užtikrinti teisių gynimą nėra tik formalus principas – tai teisinė prievolė garantuoti realų ir veiksmingą pažeistų teisių atkūrimą. Tokiu būdu konvencija įpareigoja valstybes sukurti teisinį režimą, kuriame įmanoma atsakomybė ne tik individualiam sprendėjui, bet ir sistemai kaip struktūriniam veikėjui. Todėl nacionaliniu lygmeniu bus būtina arba išplėsti, arba aiškinti civilinės atsakomybės

sampratą taip, kad ji apimtų ir vadinamąją sisteminę (beasmenę) atsakomybę, kai žala kyla ne dėl konkretaus žmogaus veiksmų, bet dėl DI sistemos konstrukcinių ar veikimo ypatumų.

Tai turi akivaizdų poveikį darbo teisės santykiams. Pavyzdžiui, darbuotojas, diskriminuotas dėl išankstinių DI sistemos modelio šališkumų, gali susidurti su įrodinėjimo sunkumais, nes nei darbdavys, nei DI tiekėjas sąmoningai nediskriminavo. Tokiu atveju, nacionalinė teisė, siekdama atitikti Vilniaus konvenciją, turės galimybę įtvirtinti atsakomybės prezumpcijos principą, kuris leistų darbuotojui inicijuoti civilinį ieškinį net esant informacijos asimetrijai, ar kai kaltės įrodyti objektyviai neįmanoma, o atsakovui – paneigti prezumpciją pateikiant įrodymus dėl tinkamos priežiūros ir prevencinių veiksmų.

Panašų reguliavimą siūlo ir Europos Parlamento tyrimų tarnyba, kuri 2020 m. parengtoje analizėje dėl civilinės atsakomybės už dirbtinį intelektą pažymėjo, kad „atsakomybės sistema negali būti grindžiama tik subjektyvumo kriterijumi, kai technologija iš esmės eliminuoja žmogaus sprendimų grandį“ (Europos Parlamento tyrimų tarnyba 2020, p. 14). Tai reiškia, kad Vilniaus konvencija veikia ne tik kaip principų deklaracija, bet ir kaip sisteminis išorinis spaudimas nacionalinės teisės transformacijai, kuri turi reaguoti į technologinę realybę, nebeapsiribodama klasikinėmis delikto ar sutarties konstrukcijomis.

Nors Vilniaus konvencija nenurodo konkrečių norminių formulių, jos įtvirtinti principai – atskaitomybė, prieinamumas, objektyvumas – neišvengiamai reikalauja, kad valstybės ne tik koreguotų turimą teisinę doktriną, bet ir kurtų novatoriškas, rizika grįstas, multiplikuotą atsakomybę leidžiančias atsakomybės sistemas, ypač tais atvejais, kai darbo santykius veikia ne žmogaus sprendimas, o sisteminė technologinė aplinka.

Nors Lietuvos teisėje egzistuoja bendrosios diskriminacijos draudimo ir neturtinės žalos kompensavimo normos, vis dar nėra aiškaus reglamentavimo, kuris leistų tiesiogiai taikyti civilinę atsakomybę už netiesioginę, sisteminę diskriminaciją, kylančią iš automatizuotų technologinių sprendimų, kai diskriminacija nėra susijusi su tiesioginiu žmogaus veiksmu. Vilniaus konvencijos 6 straipsnis tai įtvirtina kaip esminį principą: ne tik gauti sprendimą, bet ir suprasti jo motyvus bei reikšmę teisinėms pasekmėms.

Kita vertus, konvencijos 10 straipsnis įpareigoja kurti veiksmingas ir savalaikes gynimo priemones, kurios turi būti „prieinamos visiems asmenims“, nepriklausomai nuo to, ar jie gali identifikuoti sprendimo autorių. Tai itin svarbu DI kontekste, kur technologinės architektūros skaidrumas dažnai nepasiekiamas darbuotojui ar net teismui. Tokios nuostatos skatina teisės doktrinoje svarstymus dėl objektyvios atsakomybės be kaltės arba bent jau pareigos įrodyti neatsakomybę perkėlimo, ypač darbdaviui arba technologijos tiekėjui. Tokia kryptis aiškiai matoma ir ES lygiu – Europos Parlamento 2022 m. siūlyme dėl dirbtinio intelekto civilinės

atsakomybės direktyvos (Europos Komisija, 2022), kurioje siūloma grįžti prie rizika pagrįstos atsakomybės schemos.

Atsižvelgiant į šiuos pokyčius, Vilniaus konvencijos ratifikavimas Lietuvoje neišvengiamai taps testu nacionalinei civilinės atsakomybės doktrinai, reikalaujant ne tik aiškesnių normų dėl subjektinės atsakomybės (darbdavio, kūrėjo, naudotojo), bet ir funkcinio instrumentų, tokių kaip pareiga pagrįsti sprendimą, teisė į kompensaciją už netiesioginę diskriminaciją ir prevencinių priemonių taikymas. Kitaip tariant, šios konvencijos įsigaliojimas suponuoją ne kosmetinius, o struktūrinius pokyčius: tiek materialinės teisės, tiek procesinės teisės srityse.

Klasikinėje civilinėje teisėje griežtoji atsakomybė (lot. *responsabilitas sine culpa*) taikoma tais atvejais, kai žalos kilimo rizika yra struktūriškai padidinta, o jos kontrolė – neadekvati įprastiems rūpestingumo standartams. Tipiniai griežtosios atsakomybės taikymo atvejai apima tokias sritis kaip transportas, branduolinė energetika ar produktų sauga, kur veikia objektyvios rizikos režimai. Dirbtinio intelekto technologijų kontekste vis dažniau keliama analogiška pozicija: autonomiškai veikiančios sistemos, priimančios sprendimus be žmogaus tiesioginio įsikišimo, taip pat turėtų būti priskirtos prie „padidintos rizikos šaltinių“, kuriems tiktų griežtosios atsakomybės režimas.

Kaip pažymi Rimkutė (2024), Dirbtinio intelekto civilinės atsakomybės direktyva (AILD) nepagrįstai eliminuoja griežtąją atsakomybę, nors jos taikymas būtų pagrįstas siekiant mažinti įrodinėjimo našą DI padarytos žalos atveju. Tokia kritika atspindi fundamentalią šiuolaikinės teisės dilemą: kaip suderinti tradicinius atsakomybės principus, grindžiamus kaltės prezumpcija, su dirbtinio intelekto technologijų veikimo autonomiškumu ir nenumatomumu. Griežtosios atsakomybės modelis, plačiai taikomas padidintos rizikos srityse (pvz., transportas, branduolinė energetika), galėtų būti analogiškai perkeliamas į aukštos rizikos DI kontekstą, kur žalą lemia ne individualus žmogaus veiksmas, o sistemos konstrukcija ar jos taikymo architektūra.

Atsižvelgiant į tai, AILD pasirinktas modelis – remtis tik kaltės prezumpcija ir galimais procesiniais palengvinimais, neužtikrina realios teisinės apsaugos nukentėjusiajam. Kaip rodo ir praktinė darbo teisės specifika, nukentėjęs darbuotojas retai turi galimybę įrodyti algoritminės logikos neteisingumą, kūrėjo aplaidumą. Todėl darbo santykių srityje būtent objektyvios, net be kaltės grindžiamos atsakomybės formos turėtų būti plėtojamos. Rimkutės kritika šiuo atveju atlieka svarbią doktrininę funkciją, ji parodo, kad teisėkūra, net ir siekdama pažangos, rizikuoja išlaikyti neproporcingą įrodinėjimo našą silpnesniam šaliai – darbuotojui.

Europos Komisijos pasiūlymas dėl DI atsakomybės direktyvos (Europos Komisija, 2022) išreiškia aiškią intenciją judėti šia kryptimi. Dokumente nurodoma, kad, „atsižvelgiant į

technologijos autonomiškumą ir veikimo nenuspėjamumą, griežtoji atsakomybė tampa proporcinga priemone pažeistoms teisėms apginti“ (Europos Komisija, 2022, p. 6). Tai ypač aktualu darbo teisės kontekste, kai darbuotojas neturi jokių galimybių paveikti DI sprendimų logikos, tačiau patiria tiesiogines pasekmes – atleidimą, diskriminaciją ar reputacinę žalą.

Dirbtinio intelekto technologijos vis dažniau priskiriamos prie padidintos rizikos šaltinių, kuriems būdingas neprognozuojamumas, techninė autonomija ir sudėtingas priežastinis ryšys tarp sprendimo ir žalos. Kaip pažymi Militsyna (2022), tokių sistemų atveju turėtų būti taikomas privalomasis civilinės atsakomybės draudimas, siekiant užtikrinti žalos kompensavimą net ir tais atvejais, kai neįmanoma identifikuoti konkretų sprendimą priėmusio subjekto ar įrodyti kaltę (p. 156). Tai atitinka vis labiau plintančią tarptautinę praktiką, kur siekiama nebevertinti, kas kaltas, o labiau – kokia rizika kilo dėl technologijos naudojimo. Kitaip tariant, svarbiausia tampa ne tai, ar konkretus asmuo padarė klaidą, o tai, kad pati technologija – pavyzdžiui, dirbtinis intelektas – gali sukelti žalą, net jei niekas sąmoningai blogai neveikė. Todėl atsakomybė būtų priskiriama ne todėl, kad kažkas norėjo pakenkti, bet todėl, kad pati sistema gali sukelti pavojingas pasekmes. Nors Militsyna nekalba konkrečiai apie Lietuvos teisės sistemą, jos siūlymas dėl privalomo draudimo yra aktualus nacionaliniam reglamentavimui – ypač turint omenyje, kad Lietuvos CK 6.270 straipsnis leidžia objektyvią atsakomybę už padidinto pavojaus šaltinius, tačiau šiuo metu nėra aišku, ar DI sistemos gali būti priskirtos prie tokių šaltinių. Atsižvelgiant į tai, kyla poreikis normatyviai išplėsti padidintos rizikos sampratą ir apsvastyti galimybę reglamentuoti draudiminės atsakomybės režimą, kuris būtų taikomas DI kūrėjams ir tiekėjams nepriklausomai nuo jų kaltės.

Šį poreikį palaiko ir Čerka, Grigienė ir Sirbikytė (2015), siūlydami funkcinės atsakomybės modelį, kuriame kūrėjui taikoma griežtoji atsakomybė, ypač tais atvejais, kai sistemos veikimas priklauso nuo mokymosi duomenų kokybės, kurių pasekmės tampa nenumatomos. Jie teigia: „Atsakomybė be kaltės tampa būtina ten, kur atsiranda rizikos prigimtis pati savaime – tai yra, kai DI sistema veikia savarankiškai, pagal modelį, kuris neleidžia prognozuoti konkretaus sprendimo“ (Čerka et al., 2015, p. 378).

Tuo tarpu AI akto (2024) 25 straipsnyje įtvirtinta, kad tiekėjai, importuotojai ir diegėjai laikytini atsakingais, jei jie pakeičia DI sistemos funkcionalumą ar numatytąją paskirtį taip, kad ji tampa aukštos rizikos sistema. Tai kuria prielaidą griežtai objektyviai atsakomybei, kai rizikos valdymas laikomas pačių veikėjų pareiga, nepriklausomai nuo kaltės (Reglamentas (ES) 2024/1689, 25 str.).

Ši tendencija rodo, kad teisinis diskursas DI srityje evoliucionuoja nuo subjektyvios, kaltės pagrįstos atsakomybės link objektyvios, rizika grįstos schemos, kurioje svarbiausias tampa ne veiksmo neteisėtumas, o sistemos veikimo rezultato socialinis pavojingumas. Tai atveria

erdvę nacionalinės teisės transformacijai, kur atsakomybė už automatizuotą darbo santykių valdymą gali būti perkelta darbdaviui ar DI kūrėjui net nesant įprastos kaltės – pakanka, kad sistema veikė netinkamai ir sukėlė žalą. Ši pozicija tampa ypač aktuali darbo teisės kontekste, kuriame darbuotojas yra silpnesnioji šalis – informacinė asimetrija, subordinacijos santykis ir technologinis nepermatomumas lemia, kad be aiškių objektyvios atsakomybės principų darbuotojo teisės gali likti faktiškai neginamos. Todėl Vilniaus konvencijos įtvirtinti standartai turi būti įtvirtinti ne deklaratyviai, o kaip veikiančios teisės normos, keičiančios atsakomybės logiką darbo teisės santykiuose.

Taigi, dirbtinio intelekto integracija į darbo santykių reguliavimą neišvengiamai reikalauja naujų atsakomybės doktrinos kontūrų – tokių, kurie peržengtų klasikinio delikto ribas ir leistų apmąstyti atsakomybę kaip sisteminių, daugiasluoksnį reiškinį. Atsižvelgiant į dirbtinio intelekto sistemų poveikį darbo santykiams, tikslinga nacionaliniu lygmeniu persvarstyti civilinės atsakomybės reguliavimą, pripažįstant DI kaip padidinto pavojaus šaltinį. Toks kvalifikavimas pagrįstų griežtąją civilinę atsakomybę, kuriai nereikėtų nustatyti kaltės, pakaktų įrodyti, kad sistema veikė netinkamai ir sukėlė žalą darbuotojui. Be to, būtų tikslinga įtvirtinti privalomojo civilinės atsakomybės draudimo mechanizmą, kuris užtikrintų žalos atlyginimą nepriklausomai nuo to, ar ją sukėlė darbdavys, ar DI sistemos tiekėjas. Šis modelis ne tik atitiktų rizika grįstos atsakomybės logiką, bet ir leistų įtvirtinti teisinę pusiausvyrą tarp technologinės pažangos ir darbuotojų interesų apsaugos. Tokia kryptis atveria erdvę nacionalinės teisės transformacijai, kurioje atsakomybė už automatizuotą darbo santykių valdymą galėtų būti perkeliama darbdaviui ar sistemų kūrėjui vien dėl faktinio priežastinio ryšio tarp sistemos veikimo ir žalos. Darbo teisės kontekste, kuriame darbuotojas yra struktūriškai silpnesnė šalis, tokia apsauga yra ne tik racionali, bet ir būtina. Vilniaus konvencija šiuo atveju tampa svarbiu teisiniu pagrindu, kuris pabrėžia ne tik individualių teisių apsaugos būtinybę, bet ir pareigą kurti iš anksto numatytą, rizikomis grindžiamą atsakomybės sistemą. Tai reiškia, kad teisė turi gebėti atliepti situacijas, kur žalos priežastis nėra konkretaus asmens sprendimas, o sudėtingas technologinis procesas, veikiantis be tiesioginės žmogaus valios. Civilinės atsakomybės transformacija šiame kontekste tampa ne pasirinkimu, o būtinybe, grindžiama žmogaus teisių apsauga, teisingumo principu ir proporcingu rizikų paskirstymu.

IŠVADOS

1. Dirbtinio intelekto (DI) sąvoka nėra universaliai apibrėžta, tačiau teisė privalo skirti dėmesį technologinio ir norminio šios sąvokos turinio skirtumui. 2024 m. Dirbtinio intelekto aktas suformuoja pagrindinius požymius – savarankišką veikimą, autonominį duomenų apdorojimą, gebėjimą generuoti sprendimus – tačiau nacionalinėje teisėje reikalinga tiksliau apibrėžti DI sampratą, pritaikytą darbo santykių reguliavimo kontekstui. Šiuolaikinė darbo teisė privalo pripažinti, kad dirbtinis intelektas nėra vien techninė pagalbinė priemonė, tai savarankiškai galinti veikianti sistema, kuri, pasitelkdama algoritminius modelius ir žmogaus pateiktus duomenis, gali reikšmingai paveikti darbo santykių organizavimą.
2. Dirbtinio intelekto sprendimų taikymas darbo santykiuose kelia iššūkių darbuotojų pagrindinėms teisėms, ypač teisėms į privatumą, apsaugą nuo diskriminacijos ir sprendimų paaiškinamumą. Automatizuotų sistemų naudojimas darbuotojų atrankoje, veiklos vertinime ar sprendimų dėl darbo sutarčių nutraukimo priėmimo turi būti suderinamas su BDAR 22 straipsnio ir Dirbtinio intelekto akto 29 straipsnio nuostatomis. Darbdavys privalo informuoti darbuotoją apie DI sistemų taikymą, sprendimo logiką ir užtikrinti, kad sprendimai, darantys reikšmingą poveikį darbuotojo teisei padėčiai, būtų peržiūrėti žmogaus.
3. Darbdavio atsakomybė naudojant DI netampa ribota ar sąlygiška. Jis lieka atsakingas už sprendimus, priimtus jo vardu ar taikant jo integruotas technologines sistemas. Net jei sprendimą technologiškai įgyvendina trečiosios šalies sukurta DI sistema, darbdaviui kyla pareiga užtikrinti sistemos atitiktį teisės normoms, veikimo skaidrumą bei darbuotojų teisių apsaugos užtikrinimą. Atsižvelgiant į tai, nacionalinėje teisėje būtų tikslinga svarstyti Darbo kodekso papildymą nuostatomis, panašiomis į priimtas Vokietijos Betriebsverfassungsgesetz pataisose, kurios, remiantis 2021 m. Betriebsrätemodernisierungsgesetz aktu, numato darbdavio pareigą informuoti darbuotojų tarybą apie planuojamą DI diegimą (§ 90) bei reikalavimą konsultuotis dėl algoritminio vertinimo sistemų, kai jos daro įtaką darbo sąlygoms ar stebėsenai (§ 87, § 95). Tokie teisiniai instrumentai padeda išvengti darbdavių interpretacijų taikant esamas darbo teisės normas, nustato pareigą pagrįsti DI technologinių sprendimų taikymą, įgalina darbuotojų atstovus dalyvauti sprendimų priėmimo ir taip užtikrina proporcingą bei teisėtą DI integraciją į darbo santykius. Tokia teisinė aiškumo plėtra, kai įstatymo normos būtų papildomos atvejais, kurie aiškiai įvardytų konsultavimosi, informavimo procedūras būtent dėl DI užtikrintų, kad DI diegimas darbo santykiuose nesukurtų interpretacinių spragų ir būtų vykdomas laikantis skaidrumo bei darbuotojo teisių apsaugos principų.
4. Europos Sąjungos teisės aktai - BDAR ir DI aktas aiškiai nustato darbuotojo teisę į informavimą, sprendimo logiką ir žmogaus įsikišimą. Pagal BDAR 22 straipsnį, kai sprendimas

priimamas vien automatizuotai, jis negali būti galiojantis be galimybės žmogui jį peržiūrėti. Sprendimai, kurie lemia darbuotojo teisinę padėtį, pavyzdžiui, dėl priėmimo į darbą ar darbo sutarties nutraukimo negali būti grindžiami vien automatizuotomis sistemomis. Nacionalinės teisės aktuose tikslinga detaliau sureguliuoti, ar papildyti esame normas, būtent dėl DI naudojimo darbuotojų ar jų atstovų informavimo, ginčijimo ir sprendimų peržiūros procesų etapuose. Toks reglamentavimas būtent dėl DI galėtų būti įtvirtintas ir vidiniuose įmonių aktuose ir/ar kolektyvinėse sutartyse atsižvelgiant į DK normų papildymą ar naują reglamentavimą.

5. Nors klasikinė deliktinė atsakomybė pagal DK ir CK nuostatas išlieka aktuali, DI naudojimas darbo aplinkoje kelia klausimą, ar šios konstrukcijos yra pakankamos tais atvejais, kai žala darbuotojui kyla ne dėl tiesioginio žmogaus veikimo, bet dėl sisteminio, automatinio sprendimo, kurį suformuoja DI sistema. Atsižvelgiant į tai, kad DI gali būti užprogramuotas remiantis netiksliais, diskriminaciniais ar istoriškai klaidingais duomenimis, būtina svarstyti galimybę darbo teisės kontekste plėtoti rekomendacinį griežtesnės atsakomybės modelį. Tačiau tai reikala būtų papildomos doktrininės ir norminės analizės, ar DI gali būti laikomas padidinto pavojaus šaltiniu darbo teisėje. Toks modelis suteiktų darbdaviui teisę, atlyginus žalą darbuotojui, regresine tvarka reikalauti atsakomybės iš DI kūrėjo, tiekėjo ar jo draudiko, jei nustatoma, kad sprendimo pagrindas buvo sistemiškai ydinga ar diskriminacinė algoritminė logika. Šiame kontekste aktualu svarstyti privalomojo civilinės atsakomybės draudimo instituto taikymą DI kūrėjams ar tiekėjams, kaip papildomą garantinį mechanizmą, leidžiantį užtikrinti veiksmingą žalos atlyginimą darbuotojui ir paskirstyti atsakomybę visoje technologinėje tiekimo grandinėje.

6. Atsižvelgiant į darbo teisės principus, ES teisės aktus ir DI sistemų poveikį darbo organizavimui, rekomenduojama nacionaliniu lygmeniu papildyti Darbo kodeksą, siekiant užtikrinti tikslų reglamentavimą: (a) įtvirtinant darbdavio pareigą informuoti apie DI naudojimą, (b) užtikrinant darbuotojo teisę gauti paaiškinimą ir žmogaus įsikišimą, (c) aiškiai nurodant, kada darbo funkcijų pakeitimas dėl DI integravimo turi būti laikomas darbo sąlygų pakeitimu, reikalaujančiu darbuotojo sutikimo pagal DK 34 ir 45 str. Tai leistų pašalinti interpretacinius neaiškumus ir sustiprinti darbuotojo teisinę apsaugą technologinių permainų sąlygomis.

ŠALTINIŲ SĄRAŠAS

TEISĖS NORMINIAI AKTAI

Tarptautinės teisės aktai

1. Europos Taryba. (1996). *Europos socialinė chartija (pataisyta)*. Prieiga: <https://e-seimas.lrs.lt/portal/legalAct/lt/TAD/TAIS.42260> [žiūrėta 2025 m. kovo 8 d.].
2. Europos Taryba (2024) *Pagrindų konvencija dėl dirbtinio intelekto, žmogaus teisių, demokratijos ir teisinės valstybės* (CETS Nr. 225). Pasirašyta 2024 m. rugsėjo 5 d., Vilnius. Prieiga per internetą: <https://www.coe.int/en/web/artificial-intelligence/the-framework-convention-on-artificial-intelligence> [žiūrėta 2025 vasario 5d.].

Europos Sąjungos teisės aktai

3. Europos Parlamentas ir Taryba. (2002). *Direktyva 2002/14/EB dėl bendros darbuotojų informavimo ir konsultavimosi su jais sistemos sukūrimo Europos bendrijoje*. 2002 m. kovo 11 d. Europos Sąjungos oficialusis leidinys, L 80, p. 29–34. Prieiga: <https://eur-lex.europa.eu/legal-content/LT/ALL/?uri=celex%3A32002L0014> [žiūrėta 2025 m. vasario 07 d.].
4. Europos Sąjunga (2012). *Europos Sąjungos pagrindinių teisių chartija*, 2012/C 326/02, 21 str. Prieiga per: <https://eur-lex.europa.eu/legal-content/LT/TXT/?uri=CELEX%3A12012P%2FTXT> [žiūrėta 2025 vasario 3 d.].
5. Europos Parlamentas ir Taryba (2016). *Reglamentas (ES) 2016/679 (Bendrasis duomenų apsaugos reglamentas)*. Prieiga per: <https://eur-lex.europa.eu/legal-content/LT/TXT/?uri=CELEX%3A32016R0679> [žiūrėta 2025 m. kovo 15 d.].
6. Europos Parlamentas ir Taryba. (2016). *Reglamentas (ES) 2016/679 dėl fizinių asmenų apsaugos tvarkant asmens duomenis ir dėl laisvo tokių duomenų judėjimo (Bendrasis duomenų apsaugos reglamentas)*. 2016 m. balandžio 27 d. Europos Sąjungos oficialusis leidinys, L 119, p. 1–88. Prieiga: <https://eur-lex.europa.eu/eli/reg/2016/679/oj?locale=lt> [žiūrėta 2025 m. balandžio 01 d.].
7. Europos Komisija. (2021). *Pasiūlymas dėl Europos Parlamento ir Tarybos direktyvos dėl darbo skaitmeninėse platformose sąlygų gerinimo* (COM(2021) 762 final). Prieiga: <https://eur-lex.europa.eu/legal-content/LT/TXT/?uri=CELEX%3A52021PC0762> [žiūrėta 2025 m. vasario 11 d.].
8. European Commission. (2022). *Proposal for a Directive on adapting non-contractual civil liability rules to artificial intelligence (AI Liability Directive)*. COM(2022) 496 final. [Internete], Prieiga per: <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A52022PC0496> [žiūrėta 2025 kovo 20 d.].
9. Europos Parlamentas ir Taryba. (2024). *Reglamentas (ES) 2024/1084 dėl dirbtinio intelekto (Dirbtinio intelekto aktas)*. OJ L 84, 22.3.2024, p. 1–173. [interaktyvus]. Prieiga: <https://eur-lex.europa.eu/legal-content/LT/TXT/?uri=CELEX:32024R1084> [žiūrėta 2025 m. vasario 13 d.].
10. Europos Parlamentas ir Taryba (2024) *Reglamentas (ES) 2024/1689 dėl dirbtinio intelekto taisyklių (Dirbtinio intelekto aktas)*. OL L 1689, 2024-07-12. Prieiga per internetą: <http://data.europa.eu/eli/reg/2024/1689/oj> [žiūrėta 2025 m. vasario 20 d.].

11. Europos Parlamentas ir Taryba (2024). *Reglamentas (ES) 2024/1689 dėl dirbtinio intelekto – Dirbtinio intelekto aktas*. 2024 m. birželio 13 d. Europos Sąjungos oficialusis leidinys, L serija, L 2024/1689. Prieiga per internetą: <http://data.europa.eu/eli/reg/2024/1689/oj> [žiūrėta 2025 m. vasario 13 d.].

Lietuvos Respublikos teisės aktai

12. Lietuvos Respublikos Konstitucija. Valstybės žinios, 1992, Nr. 33-1014.
13. Lietuvos Respublikos civilinis kodeksas (2000), Valstybės žinios, 74-2262; 200.
14. Lietuvos Respublikos darbo kodeksas. TAR, 2016, Nr. 17373.

Soft law šaltiniai

15. Europos duomenų apsaugos valdyba (EDPB) (2017). *Gairės 3/2017 dėl duomenų apsaugos poveikio vertinimo ir jo nustatymo kriterijų pagal Reglamentą (ES) 2016/679*. Prieiga per: https://edpb.europa.eu/sites/default/files/files/file1/edpb_guidelines_2017_3_dpia_lt.pdf [žiūrėta 2025 kovo 18 d.].
16. Europos Komisija. (2019). *Patikimo dirbtinio intelekto etikos gairės*. [interaktyvus]. Prieiga: https://www.europarl.europa.eu/meetdocs/2014_2019/plmrep/COMMITTEES/JURI/DV/2019/11-06/Ethics-guidelines-AI_LT.pdf [žiūrėta 2025 m. balandžio 7 d.].
17. Europos duomenų apsaugos valdyba (2019). *Gairės 3/2019 dėl asmens duomenų tvarkymo per vaizdo įrašų sistemas*. Prieiga per: https://edpb.europa.eu/sites/default/files/files/file1/edpb_guidelines_201903_video_devices_lt.pdf [žiūrėta 2025 vasario 15 d.].
18. Europos duomenų apsaugos valdyba (EDPB) (2020). *Guidelines 05/2020 on Data Protection by Design and by Default*. Priimta 2020 m. spalio 20 d. Prieiga per: https://edpb.europa.eu/our-work-tools/our-documents/guidelines/guidelines-052020-data-protection-design-and-default_lt [žiūrėta 2025 balandžio 1 d.].
19. Europos Komisija (2020). *Baltoji knyga dėl dirbtinio intelekto. Europos požiūris į meistriškumą ir pasitikėjimą*. COM(2020) 65 final. Prieiga per: <https://eur-lex.europa.eu/legal-content/LT/TXT/?uri=CELEX%3A52020DC0065> [žiūrėta 2025 kovo 6 d.].
20. Europos duomenų apsaugos valdyba (2022). *Pareikšta pozicija dėl emocijų atpažinimo technologijų naudojimo darbo vietoje*. Prieiga per: <https://edpb.europa.eu> [žiūrėta 2025 vasario 17 d.].
21. Europos Komisija. (2025a). *Gairės dėl nepriimtinos rizikos dirbtinio intelekto praktikų*. [interaktyvus]. Prieiga: https://ec.europa.eu/commission/presscorner/detail/lt/ip_25_1013 [žiūrėta 2025 m. balandžio 10 d.].
22. Europos Komisija. (2025a). *Gairės dėl draudžiamos dirbtinio intelekto (DI) praktikos, kaip apibrėžta DI akte*. [interaktyvus]. Prieiga: <https://digital-strategy.ec.europa.eu/lt/library/commission-publishes-guidelines-prohibited-artificial-intelligence-ai-practices-defined-ai-act> [žiūrėta 2025 m. balandžio 11 d.].

23. Bagdanskis, T., Mačiulaitis, V. ir Mikalopas, M. (2018). Lietuvos Respublikos darbo kodekso komentaras. Vilnius: Registrų centras.
24. Bagdanskis, T. (2021). *Darbo sutarties nutraukimas darbdavio valia*. *Teisė*, 118, 32–46. Prieiga per: <https://www.zurnalai.vu.lt/teise/article/view/22774> [žiūrėta 2025 kovo 15 d.].
25. Bartkus, J. ir Stundys, T. (2018). *Dirbtinio intelekto teisinė atsakomybė: ko galima tikėtis?* In: Vilniaus universiteto Teisės fakulteto studentų mokslinė draugija (red.). *Teisės mokslo pavasaris 2018*. Vilnius: VĮ Registrų centras, p. 7–25. [Internete], Prieiga per: <https://www.tf.vu.lt/wp-content/uploads/2018/09/TM-pavasaris-2018-elektroninis-PDF.pdf> [žiūrėta 2025 vasario 15 d.].
26. Bitinas, A., Tartilas, J. ir Litvaitienė, J. (2011). *Socialinės apsaugos teisė*. Vilnius: Mykolo Romerio universitetas.
27. Boddington, P. (2017). *Towards a Code of Ethics for Artificial Intelligence*. Springer.
28. Breskienė, I. (2024). *Darbuotojų stebėjimas naudojant algoritminį valdymą, grįstą dirbtinio intelekto sistemomis*. *Teisė*, 128,. Prieiga: <https://www.zurnalai.vu.lt/teise/article/view/38464/37338> [žiūrėta 2025 m. vasario 26 d.].
29. Bulgakova, D. (2022). *Biometrinių duomenų apsaugos iššūkiai darbo santykiuose: subordinacijos įtaka sutikimo laisvumui*. *Teisė*, 124, p. 21–28. Prieiga per: <https://www.zurnalai.vu.lt/teise/article/view/29441> [žiūrėta 2025 balandžio 6 d.].
30. Cefaliello, A. and Kullmann, M. (2022). *Offering false security: How the draft Artificial Intelligence Act undermines fundamental workers' rights*. SSRN Working Paper No. 3993100. Prieiga per: <https://ssrn.com/abstract=3993100> [žiūrėta 2025 kovo 22 d.].
31. Čerka, P., Grigienė, J. ir Sirbikytė, D. (2015). Atsakomybės problematika informacinių technologijų taikymo sąlygomis. *Jurisprudencija*, 22(4), p. 375–391. [Interaktyvus], Prieiga per: <https://www.sciencedirect.com/science/article/pii/S026736491500062X> [žiūrėta 2025 kovo 1 d.].
32. Davulis, T. (2008). *Lietuvos darbo teisės modernizavimo perspektyvos*. *Jurisprudencija*, 1(103), pp. 22–31. Prieiga per: <https://etalpykla.lituanistika.lt/fedora/objects/LT-LDB-0001:J.04~2008~1367160611580/datastreams/DS.002.0.01.ARTIC/content> [žiūrėta 2025 kovo 13 d.].
33. Davulis, T. (2014). Lietuvos darbo teisės reforma ir Europos socialinio modelio raida. *Socialinė teorija, empirija, politika ir praktika*, 9, 67–81.
34. Davulis, T. (2017). Technologijos ir socialiniai pokyčiai darbo teisėje. Šaltinyje: *Lietuvos socialinė raida 2017*. Vilnius, Lietuvos socialinių tyrimų centras, p. 117–122. Prieiga: https://lstc.lt/wp-content/uploads/2023/02/Lietuvos_socialine_raida_2017.pdf
35. De Stefano, V. (2020). Algorithmic Bosses and What to Do About Them: Automation, Artificial Intelligence and Labour Protection. In: Mehta, P., Rao, S. S. P. (eds), *Economic and Policy Implications of Artificial Intelligence*. Springer, p. 65–86. Prieiga: https://link.springer.com/chapter/10.1007/978-3-030-45340-4_7 [žiūrėta 2025 m. kovo 5 d.].
36. De Stefano, V. ir Taes, S. (2022). *Algorithmic Management and Collective Bargaining*. *Transfer: European Review of Labour and Research*, 29(1), 21–36. Prieiga: <https://doi.org/10.1177/10242589221141055> [žiūrėta 2025 m. kovo 5 d.].

37. Floridi, L. ir Taddeo, M. (2016). *What is Data Ethics?*. *Philosophical Transactions of the Royal Society A*, 374(2083), 20160360. [ResearchGate+1](#) [Mariasaria Taddeo+1](#) [žiūrėta 2025 m. vasario 22 d.].
38. Goodfellow, I., Bengio, Y. and Courville, A. (2016). *Deep Learning*. [interaktyvus]. Prieiga per internetą: <https://www.deeplearningbook.org/> [žiūrėta 2025 m. vasario 22 d.].
39. Grumulaitis, A. (2023). Dirbtinio intelekto valdymas darbo teisėje: nauji iššūkiai ir galimi sprendimai. *Teisė*, 128, 1–20. Prieiga: <https://www.journals.vu.lt/teise/article/view/3356> [žiūrėta 2025 m. balandžio 14 d.].
40. Jordan, M. I. and Mitchell, T. M. (2015). ‘Machine learning: Trends, perspectives, and prospects’. *Science*, 349(6245), pp. 255–260. Prieiga per internetą: <https://www.science.org/doi/10.1126/science.aaa8415> [žiūrėta 2025 m. vasario 9 d.].
41. Juškevičiūtė-Vilienė, A. (2020). Dirbtinis intelektas ir konstitucinė teisė į teisingumą. *Folia Iuridica*, 93, 126–144. <https://doi.org/10.18778/0208-6069.93.08>
42. Kaplan, A. ir Haenlein, M. (2019). *Siri, Siri in my Hand: Who's the Fairest in the Land? On the Interpretations, Illustrations and Implications of Artificial Intelligence*. *Business Horizons*, 62(1), 15–25. doi:10.1016/j.bushor.2018.08.004.
43. Kellogg, K.C., Valentine, M.A. & Christin, A. (2020). *Algorithms at Work: The New Contested Terrain of Control*. *Academy of Management Annals*, 14(1), p. 366–410. <https://doi.org/10.5465/annals.2018.0174>
44. Murauskas, D. (2020). *Algoritminiai sprendimai ir teisinis reguliavimas: nuo Bendrojo duomenų apsaugos reglamento iki algoritminio valdymo*. *Teisė*, 117, p. 134–151. Prieiga: <https://www.zurnalai.vu.lt/teise/article/view/20321> [žiūrėta 2025 m. kovo 19 d.].
45. Murauskas, D. (2020). *Dirbtinis intelektas priimant teismo sprendimą – algoritmų klasifikavimas remiantis teisinio kvalifikavimo stadijomis*. *Teisė*, 115, 55–69. Prieiga per internetą: <https://doi.org/10.15388/Teise.2020.115.4> [žiūrėta 2025 kovo 17 d.].
46. Rimkutė, D. (2022). *Automatizuoto sprendimo pasekmės ir žmogaus teisės: technologijos etika ir darbo teisė*. *Teisė ir technologijos*, 5(1), p. 84–94.
47. Rimkutė, D. (2024). *Civilinės atsakomybės formavimas skaitmeniniame amžiuje: Dirbtinio intelekto atsakomybės direktyva ir Peržiūrėta atsakomybės už gaminius direktyva*. *Teisė*, 130, 129–143. DOI: <https://doi.org/10.15388/Teise.2024.130.11>
48. Scholz, B. (2020). Algorithmic Contracts. *Stanford Technology Law Review*, 23(2), 242–279. [Interaktyvus] Prieiga per: https://law.stanford.edu/wp-content/uploads/2018/03/3_SCHOLZ-FINAL_Formatted_Mar18.pdf [žiūrėta 2025 kovo 17 d.].
49. Scholz, B. (2020). *Artificial Intelligence and Responsibility: Moral and Social Considerations*. In: Smith, J. (ed.) *Ethics in the Age of AI*. Springer, pp. 245–260.
50. Todoli-Signes, Adrián (2021). Algorithmic Management: An Overview of the Research Opportunities Raised by a New Way of Working. SSRN. [Prieiga internete], Prieiga per: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3915718 [žiūrėta 2025 vasario 23 d.].
51. Vedder, A. ir Naudts, L. (2020). *Accountability for the Use of Algorithms in a Big Data Environment*. *International Review of Law, Computers & Technology*, 34(2), p. 123–143. [Internete], Prieiga per: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3407876 [žiūrėta 2025 kovo 1 d.].

52. Wachter, S., Mittelstadt, B. and Floridi, L. (2017). Why a Right to Explanation of Automated Decision-Making Does Not Exist in the General Data Protection Regulation. *International Data Privacy Law*, 7(2), 76–99. Prieiga: <https://doi.org/10.1093/idpl/ix005>
53. Wachter, S. ir Mittelstadt, B. (2019). *A Right to Reasonable Inferences: Re-Thinking Data Protection Law in the Age of Big Data and AI*. *Columbia Business Law Review*, 2019(2), p. 494–620. [Internete], Prieiga per: <https://journals.library.columbia.edu/index.php/CBLR/article/view/3424/1370> [žiūrėta 2025 vasario 2 d.].
54. Zaleskis, J. (2019). *Algoritminio sprendimų priėmimo ribos BDAR kontekste: asmens duomenų apsauga darbo teisėje*. *Teisė ir technologijos*, 6(1), 249–263. Prieiga per: <https://www.zurnalai.vu.lt/teise/article/view/38464/37338> [žiūrėta 2025 kovo 15 d.].

Elektroniniai leidiniai

55. Dastin, J. (2018). *Amazon scrapped ‘AI’ recruiting tool that showed bias against women*. Reuters, 10 October. Prieiga per: <https://www.reuters.com/article/us-amazon-com-jobs-automation-insight-idUSKCN1MK08G> [žiūrėta 2025 balandžio 8 d.].
56. EEOC (2023). *EEOC Settles Age Discrimination Lawsuit Against iTutorGroup*. Prieiga per: <https://www.eeoc.gov/newsroom/eeoc-settles-age-discrimination-lawsuit-against-itutorgroup> [žiūrėta 2025 balandžio 7 d.].
57. Europos Parlamentas. (2020). *Kas yra dirbtinis intelektas ir kaip jis naudojamas?* [interaktyvus]. Prieiga per internetą: <https://www.europarl.europa.eu/topics/lt/article/20200827STO85804/kas-yra-dirbtinis-intelektas-ir-kaip-jis-naudojamas> [žiūrėta 2025 m. vasario 14 d.].
58. IndustriAll European Trade Union. (2023). *AI at the workplace: Collective bargaining needed to ensure worker involvement*. [Internete], Prieiga per: <https://news.industriall-europe.eu/Article/924> [žiūrėta 2025 vasario 10 d.].
59. Macijauskienė, A. (2024). 6 pagrindinės dirbtinio intelekto turinio kūrimo įrankių naudojimo taisyklės darbuotojams. Prieiga per: <https://www.teise.pro/index.php/2024/01/08/6-pagrindines-dirbtinio-intelektu-turinio-kurimo-irankiu-naudojimo-taisykles-darbuotojams/>
60. McCarthy, J., Minsky, M.L., Rochester, N. and Shannon, C.E. (1955). *A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence*. Hanover: Dartmouth College. [interaktyvus]. Prieiga per internetą: <http://jmc.stanford.edu/articles/dartmouth/dartmouth.pdf> [žiūrėta 2025 m. kovo 14 d.].
61. Ministère du Travail. (2021). *La mise en place d’un dispositif de surveillance des salariés : le rôle du CSE*. [Internete], Prieiga per: <https://travail-emploi.gouv.fr/dialogue-social/le-comite-social-et-economique-cse/article/le-role-du-cse-en-cas-de-surveillance-des-salaries> [žiūrėta 2025 vasario 16 d.].
62. Tamošaitis, M. ir Gaubienė, N. (2025). *Dirbtinio intelekto atsakomybės iššūkiai. Teisė profesionaliai*, [interaktyvus]. Prieiga: <https://www.teise.pro/index.php/2025/03/20/m-tamosaitis-n-gaubiene-dirbtinio-intelektu-atsakomybes-issukiai/> [žiūrėta 2025 m. balandžio 10 d.].

Teismų praktika

63. *EEOC v. iTutorGroup, Inc.*, Civil Action No. 2:22-cv-07182, C.D. Cal. 2023. Prieiga per: <https://www.eeoc.gov/newsroom/eeoc-settles-age-discrimination-lawsuit-against-itutorgroup> [žiūrėta 2025 kovo 18 d.].
64. Europos Žmogaus Teisių Teismas (2009). *Lombardi Vallauri v. Italy*, no. 39128/05, ECHR. Prieiga per: <https://hudoc.echr.coe.int/eng?i=001-93878> [žiūrėta 2025 vasario 12 d.].
65. Europos Žmogaus Teisių Teismas (2017). *Barbulescu v. Romania*, peticijos Nr. 61496/08, Didžiosios kolegijos sprendimas 2017 m. rugsėjo 5 d. Prieiga per: <https://hudoc.echr.coe.int/eng#%7B%22itemid%22%3A%22001-177082%22%7D>
66. Europos Žmogaus Teisių Teismas (2022). *Florindo de Almeida Vasconcelos Gramaxo v. Portugal*, peticijos Nr. 26968/16, sprendimas priimtas 2022 m. gruodžio 13 d. Strasbūras: Europos Taryba. Prieiga per: [Florindo de Almeida Vasconcelos Gramaxo v. Portugal](#)
67. Lietuvos Respublikos Konstitucinis Teismas (2014). *Nutarimas dėl teisės į teisminę gynybą aiškinimo*, 2014 m. balandžio 16 d., byla Nr. 25/2011. Prieiga per: <https://www.lrkt.lt/lt/teismo-aktai/paieska/135/ta1974/content> [žiūrėta 2025 kovo 17 d.].
68. Lietuvos vyriausiasis administracinis teismas (2020). *Nutartis administracinėje byloje Nr. A-3345-822/2020*. Vilnius, 2020 m. lapkričio 12 d.
69. *Mobley v. Workday Inc.*, No. 3:23-cv-00770, N.D. Cal. 2024. Teismo dokumentas Dkt. 46, 2024 m. liepos 12 d. Prieiga per: <https://www.courtlistener.com/docket/66759458/mobley-v-workday-inc/> [žiūrėta 2025 kovo 27 d.].
70. Rechtbank Amsterdam (2021). *ECLI:NL:RBAMS:2021:5305* (Riders Union v. Uber & Ola). Amsterdamo apygardos teismas, 13 September 2021. Prieiga per: <https://uitspraken.rechtspraak.nl#!/details?id=ECLI:NL:RBAMS:2021:5305> [žiūrėta 2025 vasario 12 d.].
71. Rechtbank Den Haag (2020). *NJCM c.s. v. Staat der Nederlanden (SyRI)*, 2020 m. vasario 5 d. sprendimas, ECLI:NL:RBDHA:2020:865. Prieiga per: <https://uitspraken.rechtspraak.nl#!/details?id=ECLI:NL:RBDHA:2020:865> [žiūrėta 2025 balandžio 1 d.].
72. *Rosenbach v. Six Flags Entertainment Corp.*, 2019 IL 123186 (Illinois Supreme Court). Prieiga per: <https://courts.illinois.gov/opinions/SupremeCourt/2019/123186.pdf> [žiūrėta 2025 balandžio 1 d.].

Kiti šaltiniai

73. AI Now Institute. (2023). *Monitoring without Consent: Algorithmic Management and the Future of Work*. New York: AI Now Institute. Prieiga per: <https://ainowinstitute.org> [žiūrėta 2025 vasario 15 d.].
74. Bundesministerium der Justiz. (n.d.). §87 BetrVG – Mitbestimmungsrechte. Prieiga per internetą: https://www.gesetze-im-internet.de/betrvg/_87.html [žiūrėta 2025 vasario 11 d.].
75. Bundesrepublik Deutschland (2001). *Betriebsverfassungsgesetz (BetrVG)*, § 80. Prieiga per: https://www.gesetze-im-internet.de/betrvg/_80.html [žiūrėta 2025 kovo 7 d.].
76. Civil Rights Act of 1964, Title VII, 42 U.S.C. § 2000e-2. Prieiga per: <https://www.eeoc.gov/statutes/title-vii-civil-rights-act-1964> [žiūrėta 2025 kovo 15 d.].

77. Council of Europe (2023). *Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law (Vilnius Convention)*. Prieiga per internetą: <https://www.coe.int/en/web/artificial-intelligence/convention>
78. European Court of Auditors (2024). *Special Report: EU support for Artificial Intelligence – Ambitions and Challenges*. Luxembourg: Publications Office of the European Union.
79. European Trade Union Institute (ETUI). (2023). *AI at the workplace: Collective Rights in the Era of Algorithmic Management*. Brussels: ETUI. Prieiga per: <https://www.etui.org/publications/ai-workplace-collective-rights-era-algorithmic-management> [žiūrėta 2025 vasario 3 d.].
80. Europos Audito Rūmai (2024) *Specialioji ataskaita 08/2024: ES užmojis dirbtinio intelekto srityje – būtinas tvirtesnis valdymas ir didesnės bei tikslingesnės investicijos*. Prieiga per internetą: <https://www.eca.europa.eu/lt/publications?ref=SR-2024-08> [žiūrėta 2025 vasario 24 d.].
81. Europos Komisija. (2022). *Proposal for a Directive on adapting non-contractual civil liability rules to artificial intelligence (AI Liability Directive)*. COM(2022) 496 final.
82. Europos Parlamento tyrimų tarnyba (2020). *Dirbtinis intelektas ir civilinė atsakomybė*. Prieiga per internetą: [https://www.europarl.europa.eu/RegData/etudes/STUD/2020/621926/IPOL_STU\(2020\)621926_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/STUD/2020/621926/IPOL_STU(2020)621926_EN.pdf) [žiūrėta 2025 balandžio 3 d.].
83. Illinois General Assembly (2008). *Biometric Information Privacy Act (BIPA), 740 ILCS 14*. Prieiga per: <https://www.ilga.gov/legislation/ilcs/ilcs3.asp?ActID=3004&ChapterID=57> [žiūrėta 2025 kovo 21 d.].
84. Information Commissioner's Office (ICO) (2024). *ICO orders Serco Leisure to stop using biometric data for staff monitoring*. The Guardian, 16 April 2024. Prieiga: <https://www.theguardian.com/business/2024/apr/16/leisure-centres-scrap-biometric-systems-to-keep-tabs-on-staff-amid-uk-data-watchdog-clampdown> [žiūrėta 2025-04-19].
85. International Labour Organization (ILO). (2021). *The role of digital labour platforms in transforming the world of work*. Geneva: International Labour Office.
86. International Labour Organization (ILO). (2017). *World Employment and Social Outlook: Trends 2017*. Geneva: ILO. Prieiga per: <https://www.ilo.org/global/research/global-reports/weso/2017/lang--en/index.htm> [žiūrėta 2025 vasario 25 d.].
87. International Labour Organization. (2021). *The role of digital labour platforms in transforming the world of work*. Ženeva: International Labour Office. Prieiga per internetą: https://www.ilo.org/global/research/global-reports/weso/2021/WCMS_771749/lang--en/index.htm [žiūrėta 2025 kovo 22 d.].
88. International Labour Organization (ILO). (2022). *The Algorithmic Management of Work and its Implications in Different Contexts*. Geneva: ILO. [Internete], Prieiga per: https://www.ilo.org/wcmsp5/groups/public/---ed_emp/documents/publication/wcms_849220.pdf [žiūrėta 2025 kovo 22 d.].
89. Lietuvos socialinių tyrimų centras (LSTC). (2023). *Dirbtinio intelekto taikymo poveikis darbo santykiams Lietuvoje*. Vilnius: LSTC. Prieiga per internetą: https://www.lrs.lt/sip/getFile3?p_fid=85159 [žiūrėta 2025 kovo 15 d.].

90. Nederland (2022). *Wet op de ondernemingsraden (WOR)*, artikel 25. Prieiga per: https://wetten.overheid.nl/BWBR0002747/2022-01-01#HoofdstukVII_Artikel25 [žiūrėta 2025 vasario 18 d.].
91. République Française. (2024). *Code du travail – Article L2312-8*. Légifrance. [Internete], Prieiga per: https://www.legifrance.gouv.fr/codes/article_lc/LEGIARTI000037389486 [žiūrėta 2025 balandžio 1 d.].
92. Valstybės duomenų agentūra. (2025). Dirbtinio intelekto apibrėžimas. Oficialiosios statistikos portalas. [interaktyvus]. Prieiga: <https://osp.stat.gov.lt/> [žiūrėta 2025 m. kovo 7 d.].
93. Visuotinė lietuvių enciklopedija. (2025). *Dirbtinis intelektas*. [interaktyvus]. Prieiga: <https://www.vle.lt/straipsnis/dirbtinis-intelektas/> [žiūrėta 2025 m. vasario 2 d.].

SANTRAUKA

Dirbtinis intelektas ir naujosios technologijos: darbo teisės transformacija

Gabija Leckaitė

Skaitmeninių technologijų pažanga ir dirbtinio intelekto sistemų integracija keičia tradicinius darbo santykių modelius bei kelia naujus iššūkius civilinės ir darbo teisės atsakomybės sistemoms. Šiuolaikinės dirbtinio intelekto sistemos, veikiančios savarankiškai ir priimančios sprendimus žmogaus teisių požiūriu reikšmingose situacijose, kelia klausimą dėl atsakomybės paskirstymo tarp naudotojo – darbdavio – ir DI kūrėjo. Dėl sprendimų nepermatomumo, autonomijos ir išankstinio šališkumo rizikos, klasikinės deliktinės atsakomybės formos tampa nepakankamos.

Pirmoje darbo dalyje analizuojama atsakomybės teisinė samprata bei pagrindiniai elementai: neteisėti veiksmai, kaltė, žala ir priežastinis ryšys. Išskiriama objektyviosios ir subjektyviosios atsakomybės skirtis, aptariami alternatyvūs atsakomybės modeliai – funkcinė, inžinerinė, rizikos paskirstymo teorija – aktualūs DI kontekste.

Antroje dalyje nagrinėjama ES ir nacionalinė teisė, reglamentuojanti DI sistemų naudojimą darbo santykiuose. Aptariamos AI akto (Reglamentas (ES) 2024/1689), Dirbtinio intelekto civilinės atsakomybės direktyvos pasiūlymo (COM(2022) 496 final), taip pat Bendrojo duomenų apsaugos reglamento ir nacionalinių šaltinių normos, susijusios su atsakomybe už automatizuotus sprendimus.

Trečiojoje dalyje išskiriamos praktinės atsakomybės problemos, kai algoritminiai sprendimai sukelia žalą darbuotojui: diskriminaciją, atleidimą ar reputacinę žalą. Aptariamas Vilniaus konvencijos dėl žmogaus teisių dirbtinio intelekto srityje teisinis poveikis bei būtinybė nacionalinei doktrinai išplėsti objektyvios atsakomybės sampratą.

Darbo pabaigoje pateikiamos išvados, jog egzistuojanti atsakomybės sistema nėra pakankamai pritaikyta technologijų specifikai, todėl būtinas teisės sistemos transformavimas – tiek darbo teisės, tiek civilinės atsakomybės doktrinos, įtraukiant kolektyvinės, paskirstytos ir beasmenės atsakomybės formas.

SUMMARY

Artificial intelligence and emerging technologies: the transformation of labor law

Gabija Leckaitė

The advancement of digital technologies and the integration of artificial intelligence (AI) systems are reshaping traditional employment relations and challenging existing models of legal liability in civil and labour law. Contemporary AI systems that operate autonomously and make decisions in contexts significantly affecting human rights raise complex questions regarding the distribution of responsibility between the user – typically the employer – and the AI developer. Due to opacity, autonomy, and the risk of inherent algorithmic bias, classical tort liability models often prove insufficient.

The first part of the thesis examines the legal concept of liability and its key components: unlawful conduct, fault, damage, and causality. It distinguishes between subjective and objective liability and introduces alternative theoretical approaches, such as functional, engineering, and risk-distribution models, all of which are particularly relevant in the context of AI.

The second part explores European Union and national legislation regulating the use of AI in employment contexts. It reviews the AI Act (Regulation (EU) 2024/1689), the proposal for an AI Liability Directive (COM(2022) 496 final), the General Data Protection Regulation (GDPR), and national legal instruments governing responsibility for automated decision-making.

The third part addresses practical liability issues arising when algorithmic decisions cause harm to workers, such as discrimination, dismissal, or reputational damage. The influence of the 2023 Council of Europe Convention on AI and Human Rights (the Vilnius Convention) is analysed, emphasizing the necessity for national legal doctrines to extend the concept of objective liability to capture the structural characteristics of AI.

The thesis concludes that the current legal framework for liability does not adequately reflect the realities of technological decision-making. A transformation of legal doctrine is therefore required – both in labour law and civil liability – towards recognising collective, distributed, and impersonal forms of liability in response to the specific nature of AI-driven systems.

PRIEDAI

1 priedas. Įmonės, naudojančios dirbtinio intelekto technologijas ir naudojimo tikslai.

		Įmonės, naudojančios dirbtinio intelekto technologijas ir naudojimo tikslai proc. ¹		
		2021	2022	2023
Visos dirbtinio intelekto technologijos	Iš viso pagal darbuotojų skaičių	4,5	4,9	8,8
	10–49 darbuotojai	3,2	3,4	6,5
	50–249 darbuotojai	7,7	8,8	15,1
	250 ir daugiau darbuotojų	18,8	21,3	31,2
Rašytinė s kalbos analizė	Iš viso pagal darbuotojų skaičių	1,6	2,3	6,3
	10–49 darbuotojai	1,2	1,6	4,9
	50–249 darbuotojai	3	4,9	10,1
	250 ir daugiau darbuotojų	4	6,7	<u>23</u>

	daugiau darbuotojų			
Garsinės kalbos keitimas kompiuterio skaitomu formatu	Iš viso pagal darbuotojų skaičių	0,8	1,4	2,8
	10–49 darbuotojai	0,7	1,4	2,4
	50–249 darbuotojai	1,1	1,6	3,3
	250 ir daugiau darbuotojų	2,1	3,2	9,6
Rašytinė s ar garsinės kalbos generavimas	Iš viso pagal darbuotojų skaičių	0,9	1,3	3,6
	10–49 darbuotojai	0,8	0,9	2,8
	50–249 darbuotojai	1,1	2,5	5,6
	250 ir daugiau darbuotojų	1,7	3,5	14,6
Objektų ar asmenų atpažinimas pagal atvaizdus	Iš viso pagal darbuotojų skaičių	1,5	1,4	2,6
	10–49	0,9	0,9	1,8

	darbuotojai			
	50– 249 darbuotojai	3,1	2,6	4,6
	250 ir daugiau darbuotojų	7,9	7,7	11,8
Mašinų mokymasis duomenų analizei	Iš viso pagal darbuotojų skaičių	1,2	1,3	2,8
	10– 49 darbuotojai	0,8	0,7	2,1
	50– 249 darbuotojai	2,3	2,9	4
	250 ir daugiau darbuotojų	5	8,3	14,6
Įvairių darbo procesų automatizavimas	Iš viso pagal darbuotojų skaičių	2,3	2,4	4,4
	10– 49 darbuotojai	1,6	1,6	3,3
	50– 249 darbuotojai	3,8	4,2	6,9
	250 ir daugiau darbuotojų	11,9	13,1	<u>18,5</u>
Aplinkos	Iš	0,9	0,5	1,4

stebėjimu pagrįsti automatizuoti sprendimai sudarant sąlygas mašinų fiziniam judėjimui	viso pagal darbuotojų skaičių			
	10–49 darbuotojai	0,7	0,3	1,1
	50–249 darbuotojai	1,3	0,5	2
	250 ir daugiau darbuotojų	4	6,1	6,8

Šaltinis: Valstybės duomenų agentūra (2023), *Skaitmeninė aplinka: dirbtinis intelektas*.