



Pētījums mākslīgā intelekta sistēmas un diskriminācijas aspekti. 2. daļa.

Autore: Dr.iur. Irēna Barkāne

Saturs

Saīsinājumi	3
1. Ievads	4
2. Definīcijas.....	10
3. Tiesībaizsardzība.....	13
4. Kritiskā infrastruktūra	28
5. Migrācijas, patvēruma un robežkontroles pārvaldība	32
6. Tiesvedība un demokrātijas procesi.....	46
7. Kopsavilkums	62
Izmantotie avoti	73

Saīsinājumi

EP Konvencija - Eiropas Padomes Pamatkonvencija par mākslīgo intelektu, cilvēktiesībām, demokrātiju un tiesiskumu

ES - Eiropas Savienība

Harta – Eiropas Savienības Pamattiesību harta

MI - mākslīgais intelekts

VDAR - Vispārīgā datu aizsardzības regula

1. Ievads

Eiropas Savienībā (ES) 2024. gada 13. jūnijā tika pieņemta Eiropas Parlamenta un Padomes Regula 2024/1689, kas nosaka saskaņotas normas mākslīgā intelekta jomā – Mākslīgā intelekta akts (Mākslīgā intelekta akts).¹ Mākslīgā intelekta akts ir pirmais visaptverošais mākslīgā intelekta (MI) regulējums, kas paredz saskaņotus noteikumus par MI sistēmu laišanu tirgū, nodošanu ekspluatācijā un izmantošanu ES. Tas aizliedz vairākus nepieņemamus MI prakses veidus, kas ir pretrunā ES vērtībām un pamattiesībām, nosaka prasības augsta riska MI sistēmām, pārredzamības pienākumu attiecībā uz dažām MI sistēmām, kā arī saskaņotus noteikumus par vispārīga lietojuma MI modeļu laišanu tirgū.

Mākslīgā intelekta akts stājas spēkā 2024. gada 1. augustā. Tas kļūs tieši piemērojams visās ES dalībvalstīs, tai skaitā Latvijā, no 2026. gada 2. augusta. Tomēr atsevišķiem noteikumiem ir paredzēts atšķirīgs piemērošanas termiņš. Noteikumi par aizliegtu MI praksi un MI pratību ir jāpiemēro sešus mēnešus no Mākslīgā intelekta akta spēkā stāšanās brīža, t.i. no 2025. gada 1. februāra. Prasības vispārīga lietojuma MI modeļiem ir piemērojamas no 2025. gada 1. augusta. Tāpat noteikumi par sodiem ir piemērojami no 2025. gada 1. augusta. Augsta riska sistēmām ir noteikti atšķirīgi piemērošanas termiņi. Prasības Mākslīgā intelekta akta 6. panta 2. daļā noteiktajām augsta riska MI sistēmām, ko izmanto atsevišķās augsta riska jomās (Mākslīgā intelekta akta III pielikums), tiks piemērotas no 2026. gada 2. augusta. Savukārt prasības Mākslīgā intelekta akta 6. panta 1. punktā noteiktajām augsta riska MI sistēmām, kas ir produkta drošības sastāvdaļa vai ir produkti, uz ko attiecas daži ES saskaņošanas tiesību akti (I pielikums), ir piemērojamas gadu vēlāk, t.i. no 2027. gada 2. augusta.

Mākslīgā intelekta akta mērķis ir veicināt uz cilvēku orientēta un uzticama MI izmantošanu, vienlaikus nodrošinot augstu veselības, drošuma, ES Pamattiesību hartā² (Harta) nostiprināto pamattiesību, tostarp demokrātijas un tiesiskuma, un vides aizsardzības līmeni

¹ Eiropas Parlamenta un Padomes Regula (ES) 2024/1689 (2024. gada 13. jūnijs), ar ko nosaka saskaņotas normas mākslīgā intelekta jomā un groza Regulas (EK) Nr. 300/2008, (ES) Nr. 167/2013, (ES) Nr. 168/2013, (ES) 2018/858, (ES) 2018/1139 un (ES) 2019/.

² Eiropas Savienības Pamattiesību harta. Pieņemta 07.12.2000. OV C 2020/239, 07.06.2016.

pret MI sistēmu radītām kaitīgām sekām ES, un atbalstot inovāciju (1. panta 1. punkts). Mākslīgā intelekta aktā ir izmantota uz risku balstīta pieeja, klasificējot MI sistēmas atkarībā no dažādiem to radītajiem riska līmeņiem. MI sistēmas, kas tiek klasificētas kā augsta riska sistēmas, var radīt kaitējumu plašam pamattiesību klāstam. Mākslīgā intelekta aktā ir norādīts, ka tam, cik liela ir MI sistēmas izraisītā kaitējošā ietekme uz Hartā aizsargātajām pamattiesībām, ir īpaši svarīgi, klasificējot MI sistēmu kā augsta riska. Minētās tiesības ietver tiesības uz cilvēka cieņu, privātās un ģimenes dzīves neaizskaramību, personas datu aizsardzību, vārda un informācijas brīvību, pulcēšanās un biedrošanās brīvību, tiesības uz izglītību, patērētāju tiesību aizsardzību, darba ņēmēju tiesības, personu ar invaliditāti tiesības, dzimumu līdztiesību, intelektuālā īpašuma tiesības, tiesības uz efektīvu tiesību aizsardzību un taisnīgu tiesu, tiesības uz aizstāvību un nevainīguma prezumpciju, tiesības uz labu pārvaldību, kā arī diskriminācijas aizliegumu (Mākslīgā intelekta akta 48. apsvērums). Diskriminācijas aizliegums ir vienas no pamattiesībām, kas tiek visbiežāk pārkāptas MI sistēmu izmantošanas rezultātā, kā to parādīs Pētījumā aplūkoti prakses piemēri.

MI sistēmu klasificēšana kā augsta riska ir noteikta Mākslīgā intelekta akta 6. pantā. Pirmkārt, saskaņā ar Mākslīgā intelekta akta 6. panta 1. punktu augstu risku var radīt sistēmas, kuras ir produktu vai sistēmu drošības sastāvdaļas vai pašas ir produkti vai sistēmas, uz ko attiecas daži ES saskaņošanas tiesību akti, kas uzskaitīti I pielikumā. Tās tiek klasificētas kā augsta riska sistēmas, ja attiecīgajam produktam ir vajadzīgs trešās personas veikts atbilstības novērtējums, lai to varētu laist tirgū vai nodot ekspluatācijā saskaņā ar ES saskaņošanas tiesību aktiem. Proti, tādi produkti ir mašīnas, rotaļlietas, lifti, iekārtas un aizsardzības sistēmas, kas paredzētas lietošanai sprādzienbīstamā vidē, radioiekārtas, spiediena iekārtas, atpūtas kuģu aprīkojums, trošu ceļu iekārtas, gāzveida kurināmā iekārtas, medicīniskās ierīces, *in vitro* diagnostikas medicīniskās ierīces, pašgājēji un aviācija (Mākslīgā intelekta akta 51. apsvērums).

Otrkārt, augsta riska MI sistēmas var būt savrupas MI sistēmas, kuras nav produktu drošības sastāvdaļas vai kuras pašas nav produkti. Mākslīgā intelekta akta 6. panta 2. punkts nosaka, ka papildus 6. panta 1. punktā minētajām augsta riska MI sistēmām par augsta riska sistēmām uzskata III pielikumā minētās MI sistēmas. Tās ir klasificējamas kā augsta riska sistēmas, ja, pamatojoties uz to paredzēto nolūku, tās rada augstu kaitējuma risku personu veselībai, drošībai vai pamattiesībām, ņemot vērā gan iespējamā kaitējuma smagumu, gan

tā iespējamību, un ja tās lieto vairākās konkrēti noteiktās jomās, kā precizēts Mākslīgā intelekta aktā (Mākslīgā intelekta akta 52. apsvērums).

Mākslīgā intelekta akta III pielikumā ir uzskaitītas astoņas jomas, kurās MI izmantošana rada augstu risku: 1) biometrija, ciktāl to izmantošanu atļauj attiecīgie ES vai valsts tiesību akti; 2) kritiskā infrastruktūra; 3) izglītība un arodmācības; 4) nodarbinātība, darba ņēmēju pārvaldība un piekļuve pašnodarbinātībai; 5) piekļuve privātiem pamatpakalpojumiem un sabiedriskajiem pamatpakalpojumiem un pabalstiem un to izmantošana; 6) tiesībaizsardzība, ciktāl to izmantošanu atļauj attiecīgie ES vai valsts tiesību akti; 7) migrācijas, patvēruma un robežkontroles pārvaldība, ciktāl to izmantošanu atļauj attiecīgie ES vai valsts tiesību akti; 8) tiesvedība un demokrātijas procesi.

2024. gadā autore pēc Tiesībsarga biroja pasūtījuma izstrādāja Pētījumu “Mākslīgais intelekts un diskriminācijas aspekti”³ (turpmāk 2024. gada Pētījums), kurā tika pētīti MI sistēmu jau šobrīd radītie diskriminācijas riski četrās no astoņām Mākslīgā intelekta akta III pielikumā minētajām jomām: 1) fizisku personu biometriskā identifikācija un kategorizācija; 2) izglītība un arodapmācības; 3) nodarbinātība, darba ņēmēju pārvaldība un piekļuve pašnodarbinātībai; 4) piekļuve privātiem pamatpakalpojumiem un sabiedriskajiem pakalpojumiem un pabalstiem un to izmantošana. 2024. gada Pētījumā vispirms tika analizēts diskriminācijas aizlieguma, datu aizsardzības un MI normatīvais regulējums, kā arī vērsta uzmanība uz esošā tiesiskā regulējuma problēmām un izaicinājumiem algoritmu un MI sistēmu izmantošanas kontekstā. Pēc tam Pētījumā tika apkopoti un analizēti uzskatāmākie piemēri no starptautiskās prakses, kas saistīti ar MI sistēmu radīto diskrimināciju iepriekš minētajās četrās jomās. Pētījuma nobeigumā tika sniegts kopsavilkums, kurā apkopota būtiskākā informācija par Pētījuma rezultātiem, secinājumi un ieteikumi.

Šī pētījuma mērķis ir izpētīt pārējās četras Mākslīgā intelekta akta III pielikumā noteiktās augsta riska jomas, kas iepriekšējā 2024. gada Pētījumā netika analizētas. Konkrētāk, pētījuma mērķis ir izpētīt piemērus no Latvijas vai starptautiskās prakses par to, kā diskriminācija var izpausties šādās augsta riska jomās:

³ [Barkāne, I. \(2024\). Mākslīgais intelekts un diskriminācijas aspekti. Tiesībsarga birojs.](#)

- 1) Tiesībaizsardzība;
- 2) Kritiskā infrastruktūra;
- 3) Migrācijas, patvēruma un robežkontroles pārvaldība;
- 4) Tiesvedība un demokrātijas procesi.

Citas jomas, kurās MI sistēmas var radīt diskriminācijas riskus, ir ārpus šajā pētījumā apskatāmo jautājumu loka.

Pētījumā tiek analizēti piemēri no prakses par MI sistēmu radīto diskrimināciju šādu aizliegto kritēriju dēļ – dzimums, invaliditāte, rase, etniskā piederība, sociālais stāvoklis, vecums, seksuālā orientācija. Vienlaikus Pētījumā tiek apskatīti arī piemēri par to, kā MI sistēmu izmantošana citu aizliegto kritēriju dēļ var radīt pozitīvu/negatīvu attieksmi iepriekš minētajās jomās.

Jāņem vērā, ka atsevišķās MI izmantošanas jomās, piemēram, kritiskā infrastruktūra, tiesvedība un demokrātiskie riski, uz šo brīdi pastāv maz prakses piemēri, un tie lielākoties nav tieši saistīti ar diskriminācijas riskiem. Ņemot vērā regulējuma un prakses attīstību, Pētījumā var tikt analizēti arī cita veida piemēri, kas tieši nav saistīti ar diskrimināciju, bet var atklāt MI izmantošanas prakses problēmjautājumus saistībā ar pamattiesību aizsardzību, kā arī regulējuma attīstības un piemērošanas aspektus konkrētajās jomās.

Šajā Pētījumā tiks analizēts MI regulējums saistībā ar izpētāmajām augsta riska jomām, norādot pozitīvos un negatīvos aspektus, kā arī likuma robus, kas var traucēt pamattiesību nodrošināšanu.

2024. gada Pētījumā tika apkopta un analizēta visaptveroša informācija par diskriminācijas aizlieguma, datu aizsardzības un MI normatīvo regulējumu. ES diskriminācijas aizlieguma princips ir noteikts kā vispārīgais princips un cilvēka pamattiesības. Diskriminācijas aizliegums, kā arī sieviešu un vīriešu līdztiesība kā vispārīgie principi ir noteikti Hartas 21. pantā un 23. pantā. Diskriminācijas aizlieguma princips un tiesiskās vienlīdzības princips ir

noteikts arī Latvijas Republikas Satversmes 91. pantā⁴. Diskriminācijas aizliegumu paredz daudzi citi ES un nacionālie tiesību akti.⁵ 2024. gada Pētījums atklāja, ka neraugoties uz plašo ES diskriminācijas aizlieguma tiesiskā regulējuma darbības jomu, regulējumam, kas paredz aizsardzību pret diskrimināciju dzimuma un rases dēļ, ir daudzas nepilnības, kas ir problemātiskas algoritmiskās diskriminācijas kontekstā. Šī situācija ir vēl problemātiskāka attiecībā uz citiem aizsargātiem pamatiem — vecumu, invaliditāti, seksuālo orientāciju un reliģiju —, kuriem saskaņā ar ES tiesību aktiem ir ierobežota aizsardzība. 2024. gada Pētījums atklāja, ka datu aizsardzības regulējumam arī ir būtiska nozīme, lai novērstu algoritmu vai MI sistēmu izmantošanas rezultātā radītos diskriminācijas riskus. Piemēram, datu aizsardzības principi ir piemērojami MI sistēmām, kas apstrādā personas datus. ES Mākslīgā intelekta akts nesamazina aizsardzību, tostarp prakses aizliegumus, ko nosaka diskriminācijas novēršanas un datu aizsardzības tiesiskais regulējums, bet gan papildina esošo regulējumu.

Šajā Pētījumā minētais regulējums atkārtoti detalizēti netiks analizēts, bet tiks vērsta uzmanība uz galvenajiem problēmjautājumiem saistībā ar Pētījumā aplūkotojām MI izmantošanas augsta riska jomām.

Pētījumā ir izmantoti Latvijas, Eiropas un citu valstu autoru pētījumi par MI sistēmu un algoritmisko diskrimināciju. Pētījumā ir apkopots un analizēts ES, kā arī Latvijas tiesiskais regulējums un starptautiskā prakse.

Pētījumā ir atklāts, ka Eiropā un citur pasaulē ir sastopami prakses piemēri par MI sistēmu radītajiem diskriminācijas riskiem, kā arī citu pamattiesību aizskāruma riskiem aplūkotojās augsta riska jomās, līdzīgi kā tas tika atklāts 2024. gada Pētījumā attiecībā uz tajā analizētajām augsta riska jomām. Tajā pašā laikā Latvijā cilvēktiesību aizsardzības iestāde, Latvijas Republikas tiesībsargs, līdz šim nav izskatījis lietas, kas saistītas ar MI sistēmu radīto diskrimināciju analizētajās augsta riska jomās. Ar līdzīgām lietām nav saskārusies arī Datu valsts inspekcija kā pamattiesību aizsardzības iestāde datu aizsardzības jomā. Sagaidāms, ka līdz ar Mākslīgā intelekta akta piemērošanu un informētību par MI radītajiem riskiem

⁴ Latvijas Republikas Satversme. Pieņemta 15.02.1922. (spēkā no 07.11.1922.). *Latvijas Vēstnesis*, 01.07.1993., Nr. 43.

⁵ [Latvijas Republikas Tiesībsargs. \(2025\). Pētījums par diskriminācijas situāciju Latvijā.](#)

Eiropas uzraudzības iestādes arvien vairāk izskatīs lietas par MI sistēmu radītajiem riskiem. Pētījumā veiktā regulējuma un prakses analīze sniedz ieskatu, kādiem aspektiem jāpievērš īpaša uzmanība, lai nodrošinātu MI sistēmu atbilstību tiesiskajam regulējumam un atbildīgu izmantošanu.

Pētījuma nobeigumā ir sniegts kopsavilkums, kurā apkopoti pētījuma rezultāti, secinājumi un ieteikumi. Pētījumā ir ieteikts efektīvi nodrošināt atbilstību Mākslīgā intelekta aktā un citos tiesību aktos noteiktajām tiesiskajām prasībām, lai mazinātu aizspriedumus un ievērotu pamattiesības. Nodrošinot atbilstību tiesiskajām prasībām (piemēram, ieviešot riska pārvaldības sistēmu, ievērojot pārredzamības un cilvēka uzraudzības prasības) un veicinot izpratni par aizspriedumu atklāšanas un mazināšanas stratēģijām, tiesībaizsardzības iestādes var atbildīgi izmantot MI sistēmas savā darbībā.

2. Definīcijas

Aizspriedums – uzskats (maldīgs, bieži vien novecojis) kas izraisa netaisnīgu spriedumu vai nepareizu rīcību pret personu, grupu vai ideju. Aizspriedumi jeb neobjektivitāte var rasties no datiem, ko izmanto MI modeļu apmācībai, vai no pašu algoritmu dizaina.⁶

Aizliegtie kritēriji – īpašības, iezīmes vai rakstura iezīmes, pamatojoties uz kurām saskaņā ar tiesību aktiem personas diskriminācija ir aizliegta, piemēram, rase, ādas krāsa, tautība un etniskā izcelsme, valoda, dzimšana un izcelsme, dzimums, vecums, invaliditāte, ģenētiskās īpašības, seksuālā orientācija, reliģiskā pārliecība, politiskā un cita pārliecība, pasaules uzskats, partijas piederība, sociālais stāvoklis un sociālā izcelsme, dienesta stāvoklis, manta.

Diskriminācija – salīdzināmā situācijā mazāk labvēlīga attieksme pret indivīdiem, pamatojoties uz aizliegtu kritēriju.

Biometriskie dati – personas dati pēc specifiskas tehniskas apstrādes, kuri attiecas uz fiziskas personas fiziskajām, fizioloģiskajām vai uzvedības pazīmēm, piemēram, sejas attēli vai daktiloskopiskie dati.

Biometriskās kategorizācijas sistēma – MI sistēma, kuras mērķis ir noteikt fizisku personu piederību konkrētām kategorijām, pamatojoties uz viņu biometriskajiem datiem, ja vien tas nepapildina citus komercpakalpojumus un nav absolūti nepieciešams objektīvu tehnisku iemeslu dēļ.

Biometriskās tālidentifikācijas sistēma – ir MI sistēma, kuras nolūks ir identificēt fiziskas personas bez to aktīvas iesaistīšanas, parasti no attāluma salīdzinot personas biometriskos datus ar atsauces datubāzē esošajiem biometriskajiem datiem.

Dzīlviltojums – MI ģenerēts vai manipulēts attēls, audio vai video saturs, kurš atgādina reālas personas, priekšmetus, vietas, vienības vai notikumus un maldinātu personu, ka attiecīgais saturs ir autentisks vai patiess.

⁶ Šī un pārējās sadaļā sniegtās definīcijas noteiktas Mākslīgā intelekta akta 3. pantā, tajā norādītajās ES tiesību normās un pētījumā [Europol \(2025\). AI bias in law enforcement - A practical guide, Europol Innovation Lab observatory report.](#)

MI neobjektivitāte – sistemātiskas kļūdas vai aizspriedumi MI sistēmas datos, algoritmos vai rezultātos, kas negodīgi dod priekšroku vai nostāda neizdevīgā stāvoklī noteiktas grupas vai indivīdus.

MI sistēma – mašinizēta sistēma, kura projektēta darboties ar dažādiem autonomijas līmeņiem, kura var pēc ieviešanas būt adaptīva, un kura eksplīcītiem vai implīcītiem mērķiem secina no informācijas, ko tā saņem, kā ģenerēt iznākumus, piemēram, prognozes, saturu, ieteikumus vai lēmumus, kas var ietekmēt fizisko vai virtuālo vidi.

MI pratība – prasmes, zināšanas un izpratne, kas nodrošinātājiem, uzturētājiem un skartajām personām dod iespēju, ņemot vērā to tiesības un pienākumus šīs regulas kontekstā, veikt MI sistēmu informētu uzturēšanu, kā arī gūt izpratni par MI iespējām un riskiem un par potenciālo kaitējumu, ko tas var radīt.

Nodrošinātājs – fiziska vai juridiska persona, publiska iestāde, aģentūra vai cita struktūra, kura par samaksu vai par brīvu izstrādā MI sistēmu vai vispārīga lietojuma MI modeli vai kura liek izstrādāt MI sistēmu vai vispārīga lietojuma MI modeli un laiž to tirgū vai nodod MI sistēmu ekspluatācijā ar savu nosaukumu/vārdu vai preču zīmi.

Personas dati – jebkura informācija, kas attiecas uz identificētu vai identificējamu fizisku personu ("datu subjekts"); identificējama fiziska persona ir tāda, kuru var tieši vai netieši identificēt, jo īpaši atsaucoties uz identifikatoru, piemēram, minētās personas vārdu, uzvārdu, identifikācijas numuru, atrašanās vietas datiem, tiešsaistes identifikatoru vai vienu vai vairākiem minētajai fiziskajai personai raksturīgiem fiziskās, fizioloģiskās, ģenētiskās, garīgās, ekonomiskās, kultūras vai sociālās identitātes faktoriem.

Profilēšana – jebkura veida automatizēta personas datu apstrāde, kas izpaužas kā personas datu izmantošana nolūkā izvērtēt konkrētus ar fizisku personu saistītus personiskus aspektus, jo īpaši analizēt vai prognozēt aspektus saistībā ar minētās fiziskās personas sniegumu darbā, ekonomisko situāciju, veselību, personīgām vēlmēm, interesēm, uzticamību, uzvedību, atrašanās vietu vai pārvietošanos.

Reāllaika biometriskās tālidentifikācijas sistēma – biometriskās tālidentifikācijas sistēma, kurā biometrisko datu uztveršana, salīdzināšana un identificēšana notiek bez būtiskas

aiztures, ietverot ne tikai tūlītēju identifikāciju, bet arī ierobežotu īslaicīgu aizturi, lai novērstu apiešanu.

Tiesībaizsardzība – darbības, ko veic tiesībaizsardzības iestādes vai to vārdā, lai novērstu, izmeklētu, atklātu noziedzīgus nodarījumus vai sauktu pie atbildības par tiem, vai izpildītu kriminālsodus, tostarp lai pasargātu no draudiem sabiedriskajai drošībai un tos novērstu.

Tiesībaizsardzības iestāde – publiska iestāde, kuras kompetencē ir novērst, izmeklēt, atklāt noziedzīgus nodarījumus vai saukt pie atbildības par tiem, vai izpildīt kriminālsodus, tostarp pasargāt no draudiem sabiedriskajai drošībai un tos novērst; vai cita struktūra vai vienība, kurai dalībvalsts tiesībās uzticēts īstenot publisku varu un publiskas pilnvaras ar mērķi novērst, izmeklēt atklāt noziedzīgus nodarījumus vai saukt pie atbildības par tiem, vai izpildīt kriminālsodus, tostarp pasargāt no draudiem sabiedriskajai drošībai un tos novērst.

Uzturētājs – fiziska vai juridiska persona, publiskā sektora iestāde, aģentūra vai cita struktūra, kura lieto MI sistēmu, kas ir tās pārziņā, izņemot gadījumus, kad MI sistēmu lieto personiskas neprofesionālas darbības veikšanai.

3. Tiesībaizsardzība

Saskaņā ar Mākslīgā intelekta aktu par augsta riska MI sistēmām ir uzskatāmas sistēmas, ko izmanto tiesībaizsardzības jomā, ciktāl to izmantošanu atļauj attiecīgie ES vai valsts tiesību akti (6. panta 2. punkts, III pielikuma 6. punkts). Augstu risku rada piecu veidu MI sistēmas, ko paredzēts izmantot tiesībaizsardzības iestādēs vai to vārdā, vai ES iestādēs, struktūrās, birojos vai aģentūrās, kas atbalsta tiesībaizsardzības iestādes, vai to vārdā:

- 1) lai novērtētu risku, ka fiziska persona varētu kļūt par cietušo personu noziedzīgos nodarījumos;
- 2) kas atbalsta tiesībaizsardzības iestādes, piemēram, melu detektoru un līdzīgi rīki;
- 3) pierādījumu uzticamības izvērtēšanai, noziedzīgu nodarījumu izmeklēšanas vai kriminālvajāšanas gaitā;
- 4) lai novērtētu risku, ka fiziska persona varētu izdarīt vai atkārtoti izdarīt nodarījumu, pamatojoties ne tikai uz fizisku personu profilēšanu, kas minēta Direktīvas (ES) 2016/680 (Policijas direktīvas)⁷ 3. panta 4. punktā, vai lai novērtētu fizisku personu vai grupu personības iezīmes un rakstura īpašības vai agrāku noziedzīgu rīcību;
- 5) Policijas direktīvas 3. panta 4. punktā minētajai fizisku personu profilēšanai noziedzīgu nodarījumu atklāšanas, izmeklēšanas vai kriminālvajāšanas gaitā.

Attiecībā uz pēdējiem diviem gadījumiem, Policijas direktīvas 3. panta 4. punkts nosaka, ka profilēšana ir “jebkura veida automatizēta personas datu apstrāde, kas izpaužas kā personas datu izmantošana nolūkā izvērtēt konkrētus ar fizisku personu saistītus personiskus aspektus, jo īpaši analizēt vai prognozēt aspektus saistībā ar minētās fiziskās personas sniegumu darbā, ekonomisko situāciju, veselību, personīgām vēlmēm, interesēm, uzticamību, uzvedību, atrašanās vietu vai pārvietošanos”.

Mākslīgā intelekta aktā ir izdalītas vairākas cita veida augsta riska MI sistēmas, ko arī izmanto tiesībaizsardzības iestādes. Mākslīgā intelekta akta III pielikuma 7. punktā kā augsta riska joma ir noteikta migrācijas, patvēruma un robežkontroles pārvaldība, bet 8. punktā –

⁷ [Eiropas Parlamenta un Padomes Direktīva \(ES\) 2016/680 \(2016. gada 27. aprīlis\) par fizisku personu aizsardzību attiecībā uz personas datu apstrādi, ko veic kompetentās iestādes \[...\]. OVL 119, 04.05.2029](#)

tiesvedības joma. Šīs jomas tiks aplūkotas darba turpinājumā. Saskaņā ar Mākslīgā intelekta akta III pielikuma 1. punktu par augsta riska MI sistēmām ir uzskatāmas arī MI sistēmas, ko izmanto biometrijā (biometriskās tālidentifikācijas sistēmas; biometriskās kategorizācijas sistēmas; emociju atpazīšanas sistēmas), kas tika aplūkota 2024. gada Pētījumā un šajā pētījumā netiks analizētas.

Jāņem vērā, ka saskaņā ar Mākslīgā intelekta akta III pielikuma 6. punktu augstu risku rada MI sistēmas tiesībaizsardzības jomā “ciktāl to izmantošanu atļauj attiecīgie ES vai valsts tiesību akti”. Šai norādei ir būtiska nozīme, jo MI sistēmas tiesībaizsardzības jomā var tikt uzskatītas arī par aizliegtām saskaņā ar citiem tiesību aktiem. Mākslīgā intelekta akta preambulā ir uzsvērts, ka fakts, ka MI sistēma saskaņā ar Mākslīgā intelekta aktu ir klasificēta kā augsta riska MI sistēma, nebūtu jāinterpretē kā norāde, ka sistēmas lietošana ir likumīga citu ES tiesību aktu vai ar tiem saderīgu valsts tiesību aktu ietvaros, piemēram, attiecībā uz personas datu aizsardzību, melu detektoru un tamlīdzīgu rīku vai citu sistēmu lietošanu fizisku personu emocionālā stāvokļa noteikšanai. Šāda lietošana būtu jāturpina tikai saskaņā ar piemērojamajām prasībām, kas izriet no Hartas un no attiecīgajiem ES un dalībvalstu tiesību aktiem. Nebūtu jāuzskata, ka Mākslīgā intelekta akts nodrošina juridisko pamatu personas datu, tostarp – attiecīgā gadījumā – īpašu kategoriju personas datu, apstrādei, ja vien Mākslīgā intelekta aktā nav konkrēti noteikts citādi (Mākslīgā intelekta akta 63. apsvērums). Mākslīgā intelekta akts nesamazina aizsardzību, tostarp prakses aizliegumus, ko nosaka diskriminācijas novēršanas un datu aizsardzības tiesiskais regulējums, bet gan papildina esošo regulējumu.

Mākslīgā intelekta akta 2. pantā ir paredzēti vairāki vispārīgi izņēmumi no Mākslīgā intelekta akta darbības jomas, kas ir būtiski, lai pilnībā izprastu tā III pielikumā uzskaitīto augsta riska jomu praktisko piemērošanu. Saskaņā ar 2. panta 3. punktu Mākslīgā intelekta akts neattiecas uz MI sistēmām, ja tās tiek laistas tirgū, nodotas ekspluatācijā vai izmantotas vienīgi militāros, aizsardzības vai valsts drošības nolūkos. Mākslīgā intelekta akta 24. apsvērumā ir sīkāk paskaidrots, kā būtu jāinterpretē jēdziens “vienīgi”, un kad MI sistēma, ko izmanto šādiem nolūkiem, tomēr var būt Mākslīgā intelekta akta darbības jomā. Piemēram, ja MI sistēma, kas laista tirgū, nodota ekspluatācijā vai tiek lietota militārām, aizsardzības vai valsts drošības vajadzībām, (kādu laiku vai pastāvīgi) izmanto citiem

mērķiem, piemēram, civiliem vai humāniem mērķiem, tiesībaizsardzības vai sabiedriskās drošības mērķiem, šāda sistēma ir Mākslīgā intelekta akta darbības jomā. Tādā gadījumā iestādei, kas izmanto MI sistēmu citiem nolūkiem, būtu jānodrošina MI sistēmas atbilstība Mākslīgā intelekta aktam, ja vien sistēma jau neatbilst minētajam aktam, – pirms tādas izmantošanas tas ir jāpārbauda. Tomēr tam, ka MI sistēma var būt Mākslīgā intelekta akta darbības jomā, nebūtu jāietekmē valsts drošības, aizsardzības un militārās darbības veicošo iestāžu spēja izmantot minēto sistēmu valsts drošības, militārām un aizsardzības vajadzībām neatkarīgi no tā, kāda veida iestāde veic minētās darbības. Piemēram, ja valsts izlūkošanas aģentūra ir uzdevusi valsts drošības aģentūrai vai privātam operatoram izmantot reāllaika biometriskās tālidentifikācijas sistēmas valsts drošības vajadzībām (piemēram, lai vāktu izlūkdatumus), šāda izmantošana būtu izslēgta no Mākslīgā intelekta akta darbības jomas.⁸ Skaidra valsts drošības izslēgšana no akta darbības jomas ir īpaši svarīga, ja MI sistēmas tiek laistas tirgū, nodotas ekspluatācijā vai lietotas tiesībaizsardzības vajadzībām, kas ir Mākslīgā intelekta akta darbības jomā.⁹

Eiropas Komisija 2025. gada 29. jūlijā pieņemtajās Pamatnostādnēs par aizliegtu mākslīgā intelekta praksi, kas noteikta Mākslīgā intelekta aktā (turpmāk – EK Pamatnostādnes)¹⁰ skaidro, ka MI prakse, kas aizliegta saskaņā ar Mākslīgā intelekta akta 5. pantu, būtu jāņem vērā saistībā ar MI sistēmām, kas klasificētas kā augsta riska sistēmas saskaņā ar Mākslīgā intelekta akta 6. pantu, jo īpaši tām, kas uzskaitītas III pielikumā. Tas ir tāpēc, ka dažos gadījumos tādu MI sistēmu izmantošanu, kas klasificētas kā augsta riska sistēmas, konkrētos gadījumos var uzskatīt par aizliegtu praksi, ja ir izpildīti visi viena vai vairāku Mākslīgā intelekta akta 5. pantā paredzēto aizliegumu nosacījumi. Savukārt lielākā daļa MI sistēmu, kurām nepiemēro Mākslīgā intelekta akta 5. pantā uzskaitīto aizliegumu, tiks klasificētas kā augsta riska sistēmas.

EK Pamatnostādnēs ir arī vērsta uzmanība, ka MI sistēmas, kas izņēmuma kārtā netiek uzskatītas par augsta riska sistēmām, pamatojoties uz Mākslīgā intelekta akta 6. panta 3. punktu, lai gan uz tām attiecas III pielikumā minētais augsta riska lietošanas gadījums,

⁸ [Eiropas Komisija. \(2025\). Komisijas paziņojums. Komisijas Pamatnostādnes par aizliegtu mākslīgā intelekta praksi, kas noteikta Regulā \(ES\) 2024/1690 \(MI aktā\).](#)

⁹ Turpat.

¹⁰ Turpat.

joprojām var būt Mākslīgā intelekta akta 5. pantā noteikto aizliegumu darbības jomā. Mākslīgā intelekta akta 6. panta 3. punkts nosaka, ka III pielikumā minēto MI sistēmu neuzskata par augsta riska sistēmu, ja tā nerada būtisku kaitējuma risku fizisku personu veselībai, drošībai vai pamattiesībām, tostarp būtiski neietekmē lēmumu pieņemšanas iznākumu. Minētais izņēmums ir piemērojams, ja ir izpildīts jebkurš no šādiem nosacījumiem: a) MI sistēma ir paredzēta šaura procedūras uzdevuma veikšanai; b) MI sistēma ir paredzēta iepriekš pabeigtas cilvēka darbības rezultāta uzlabošanai; c) MI sistēma ir paredzēta lēmumu pieņemšanas modeļu vai noviržu no iepriekšējiem lēmumu pieņemšanas modeļiem atklāšanai, un tā nav paredzēta, lai aizstātu vai ietekmētu iepriekš pabeigtu cilvēka veiktu novērtējumu, bez pienācīgas cilvēka veiktas pārskatīšanas; vai d) MI sistēma ir paredzēta sagatavošanās uzdevuma veikšanai saistībā ar novērtējumu, kas ir būtisks III pielikumā uzskaitīto lietošanas gadījumu nolūkos. Minētās normas pēdējais teikums nosaka, ka Mākslīgā intelekta akta III pielikumā minētu MI sistēmu vienmēr uzskata par augsta riska MI sistēmu, ja MI sistēma veic fizisku personu profilēšanu.

Mākslīgā intelekta akta 5. panta pirmā daļa paredz astoņus aizliegtas MI prakses veidus:

- 1) kaitējošas manipulācijas un maldināšana;
- 2) kaitējoša neaizsargātības izmantošana;
- 3) sociālais novērtējums;
- 4) atsevišķu noziedzīgu nodarījumu riska novērtējums un prognozēšana;
- 5) nemērķorientēta rasmošana sejas atpazīšanas datubāzu izstrādei;
- 6) emociju izsecināšana;
- 7) biometrisku datu kategorizācija;
- 8) reāllaika biometriskā tālidentifikācija.

Ir svarīgi nošķirt augsta riska MI sistēmas no aizliegtas MI prakses. Piemēram, MI sistēmu izmantošana biometrijas jomā var tikt klasificēta kā aizliegta MI prakse. Saskaņā ar Mākslīgā intelekta akta 5. panta 1. punkta h) apakšpunktu ir aizliegta MI sistēmu lietošana fizisku personu reāllaika biometriskajai tālidentifikācijai publiski pieklūstamās vietās tiesībaizsardzības nolūkos, izņemot izsmēloši uzskaitītas un šauri definētas situācijas, kurās sistēmu izmantošana ir absolūti nepieciešama būtisku sabiedrības interešu īstenošanai. MI sistēmas, ko izmanto biometrijas jomā, var izmantot arī citās augsta riska jomās, piemēram,

tiesībaizsardzības, migrācijas, patvēruma un robežkontroles pārvaldības, kā arī tiesvedības un demokrātisko procesu jomā. Kā minēts iepriekš, MI sistēmas, ko izmanto biometrijas jomā, tika analizētas 2024. gada Pētījumā.

Vēl viens gadījums, kad jānošķir augsta riska un aizliegta MI prakse, ir MI izmantošana noziedzīgu nodarījumu riska novērtēšanā un prognozēšanā jeb prognozējošai policijas darbībai. Saskaņā ar Mākslīgā intelekta akta 5. panta 1. punkta d) apakšpunktu ir aizliegtas MI sistēmas, kas novērtē vai prognozē risku, ka fiziska persona varētu izdarīt noziedzīgu nodarījumu, pamatojoties tikai uz profilēšanu vai personiskajām īpašībām un rakstura iezīmēm. Minētā norma paredz, ka aizliegums neattiecas uz MI sistēmu, kuru izmanto, lai atbalstītu cilvēka veiktu novērtējumu par personas iesaistīšanos noziedzīgā darbībā, kurš jau ir balstīts uz objektīviem un pārbaudāmiem faktiem, kas ir tieši saistīti ar šādu darbību.

Saskaņā ar Mākslīgā intelekta akta III pielikuma 6. punkta d) apakšpunktu MI sistēmas, uz kurām neattiecas aizliegums un kuras paredzēts izmantot tiesībaizsardzības iestādēs vai to vārdā, vai ES iestādēs, struktūrās, birojos vai aģentūrās tiesībaizsardzības iestāžu atbalstam, lai novērtētu risku, ka fiziska persona pirmo reizi vai atkārtoti izdara pārkāpumu, ne tikai pamatojoties uz profilēšanu vai personisko īpašību un rakstura iezīmju novērtēšanu, vai iepriekšēju noziedzīgu uzvedību, klasificē kā “augsta riska” MI sistēmas. Šādām MI sistēmām ir jāatbilst visām prasībām un pienākumiem, kas noteikti Mākslīgā intelekta aktā, tai skaitā cilvēka virsvadības prasībai (Mākslīgā intelekta akta 14. un 26. pants).

EK Pamatnostādnes¹¹ skaidro, ka cilvēka virsvadība ir jāuztic personām ar nepieciešamo kompetenci, apmācību un pilnvarām, kurām būtu jāspēj pienācīgi izprast MI sistēmas spējas un ierobežojumi, pareizi interpretēt tās iznākumu un novērst automatizācijas neobjektivitātes risku. Šīm personām vajadzētu būt skaidrām procedūrām, apmācībai un nepieciešamajai kompetencei un pilnvarām, lai jēgpilni novērtētu MI sistēmas iznākumus. Šajā konkrētajā gadījumā cilvēka veiktam novērtējumam būtu jānodrošina, ka jebkura MI prognoze vai novērtējums par personas, kas izdara noziedzīgu nodarījumu, risku ir balstīts uz objektīviem un pārbaudāmiem faktiem, kas saistīti ar noziedzīgu darbību. Minētajām

¹¹ [Eiropas Komisija. \(2025\). Komisijas paziņojums. Komisijas Pamatnostādnes par aizliegtu mākslīgā intelekta praksi, kas noteikta Regulā \(ES\) 2024/1690 \(MI aktā\).](#)

personām būtu arī jāiejaucas, lai izvairītos no negatīvām sekām vai riskiem vai apturētu MI sistēmas izmantošanu, ja tā nedarbojas, kā paredzēts.

EK Pamatnostādnēs tiek skaidrotas Mākslīgā intelekta akta III pielikuma 6. punkta d) apakšpunktā noteiktās augsta riska MI sistēmas, norādot, ka riska novērtējumus, lai novērtētu vai prognozētu noziedzīga nodarījuma izdarīšanas risku, bieži dēvē par individuālu noziedzīga nodarījuma prognozēšanu vai noziedzīga nodarījuma paredzēšanu. Lai gan nav vispārpieņemtas noziedzīga nodarījuma prognozēšanas vai noziedzīga nodarījuma paredzēšanas definīcijas, šie termini kopumā attiecas uz dažādām progresīvām MI tehnoloģijām un analītiskajām metodēm, ko piemēro lielam skaitam bieži vien vēsturisku datu (tai skaitā sociālekonomiskajiem datiem, kā arī policijas reģistru datiem u. c.), kurus kopā ar kriminoloģijas teorijām izmanto, lai prognozētu noziedzību kā pamatu policijas un tiesībsardzības stratēģijām un rīcībai, lai apkarotu, kontrolētu un novērstu noziedzību. Noziedzīgu nodarījumu prognozēšanas MI sistēmas identificē modeļus vēsturiskos datos, sasaistot rādītājus ar noziedzīga nodarījuma iespējamību un pēc tam kā prognozējošu iznākumu ģenerējot riska novērtējumu. Piemēram, šādas sistēmas var izmantot, lai plānotu policijas darba grupas, uzraudzītu augsta riska situācijas un veiktu to personu kontroli, kuras saskaņā ar prognozēm ir iespējami (atkārtoti) likumpārkāpēji. Tomēr šāda vēsturisko datu izmantošana par veiktajiem noziegumiem, lai prognozētu personu turpmāko uzvedību, var turpināt vai pat pastiprināt neobjektivitāti. Šāda datu izmantošana var arī izraisīt to, ka būtiski individuālie apstākļi netiek ņemti vērā, ja šie apstākļi nav daļa no datu kopas vai netiek ņemti vērā algoritmos, ar kuriem darbojas konkrētā MI sistēma.¹²

Eiropola 2025. gada pētījumā ir vērsta uzmanība uz neobjektivitātes risku, ko rada MI sistēmu izmantošana tiesībsardzības jomā. Lai gan MI tehnoloģiju pielietojumi tiesībsardzības iestādēs, piemēram, paredzošā policijas darbība, automatizēta modeļu identificēšana un uzlabota datu analīze, sniedz ievērojamas priekšrocības, tām ir arī raksturīgi riski. Šie riski rodas no aizspriedumiem, kas iestrādāti MI sistēmu projektēšanā, izstrādē un ieviešanā, kas var veicināt diskrimināciju, pastiprināt sabiedrības nevienlīdzību un apdraudēt tiesībsardzības darbību integritāti. MI radītās neobjektivitātes var rasties

¹² [Eiropas Komisija. \(2025\). Komisijas paziņojums. Komisijas Pamatnostādnēs par aizliegtu mākslīgā intelekta praksi, kas noteikta Regulā \(ES\) 2024/1690 \(MI aktā\).](#)

jebkurā MI sistēmas dzīves cikla posmā. Projektēšanas fāzē vēsturiskās un reprezentācijas neobjektivitātes bieži atspoguļo nevienlīdzīgus sabiedrības modeļus, kas noved pie sagrozītiem rezultātiem. Izvietošanas fāzē ļaunprātīgas izmantošanas neobjektivitāte un pārmērīga paļaušanās uz MI rezultātiem (automatizācijas neobjektivitāte) var novest pie nepareizām darbībām, piemēram, nevajadzīgas novērošanas vai nepareiziem arestiem. Šīs problēmas ir īpaši izteiktas prognozējošajā policijas darbā, kur neobjektīvi dati un atgriezeniskās saites cilpas var pastiprināt stereotipus un nesamērīgi vērsties pret marginalizētām kopienām.¹³

Turpmāk tiks aplūkoti konkrēti prakses piemēri, kas atklāj, kā MI sistēmas tiesībaizsardzības jomā var radīt diskriminācijas risku.

Prakses piemēri

VioGén sistēma vardarbības riska ģimenē novērtēšanai

Sistēmas, kas izmantotas, lai novērtētu risku, ka fiziska persona varētu kļūt par noziedzīgos nodarījumos cietušo personu, noteiktas kā augsta riska sistēmas saskaņā ar Mākslīgā intelekta akta III pielikuma 6. punkta a) apakšpunktu. Viens no piemēriem, kas parāda šādu sistēmu diskriminācijas risku, ir 2007. gadā Spānijas Iekšlietu ministrijas ieviestā automatizētā sistēma VioGén, kas tika paredzēta kā riska novērtēšanas rīks vardarbības ģimenē riska līmeņu atpazīšanai.¹⁴ Sistēma tika ieviesta, lai aizsargātu sievietes un bērnus no vardarbības ģimenē un novērtētu sievietēm draudošās agresijas riska pakāpi, un piešķirtu vērtējumu, kas noteiktu policijas aizsardzības līmeni, kas viņām būtu jāsaņem.

Lai gan sistēmas mērķis liekas atbalstāms, tās ieviešanas laikā radās vairākas problēmas. 95% gadījumu tika piešķirts automātisks riska vērtējums un tikai 3% gadījumi tika novērtēti kā "vidēja" vai "augsta" riska, kas tika noteikts kā minimālais sliekšnis, lai policija iejauktos.

¹³ [Europol \(2025\). AI bias in law enforcement - A practical guide, Europol Innovation Lab observatory report.](#)

¹⁴ [Sobrinó-García, I. \(2021\). Artificial Intelligence Risks and Challenges in the Spanish Public Administration: An Exploratory Analysis through Expert Judgements. Administrative Sciences. 11\(3\);](#) [Martínez Garay, L., Boix Palop, A., Briz Redón, Á. et.al. \(2022\). Three predictive policing approaches in Spain: Viogén, RisCanvi and Veripol. Assessment from a human rights perspective;](#) [Orakhelashvili, A. \(2024\). Real World Cases of Annex III AI Act applications that posed risks to fundamental rights.](#)

Iespējamās sekas sistēmas izmantošanai varētu būt, ka laika posmā no 2003. līdz 2021. gadam 71 sieviete, kas bija iesniegusi ziņojumu un nesaņēma aizsardzību, tika nogalināta.

Pētnieki arī vērš uzmanību uz iespējamu diskrimināciju **tautības vai izcelsmes** valsts dēļ. VioGén ārvalstīs dzimušajām sievietēm piešķīra zemākus riska līmeņus nekā Spānijā dzimušajām sievietēm, izlīdzinot to atbilstoši iedzīvotāju skaitam, ko varētu uzskatīt par netiešu diskrimināciju.¹⁵

Līdzās sistēmas uzticamībai, tika konstatēta arī pārredzamības problēma. Tikai 35% sieviešu, kuras izmantoja VioGén sistēmu, zināja savu vērtējumu. Tas liecina, ka sistēma nebija pārskatāma un sievietēm, kuras uz to paļāvās, netika sniegta atgriezeniskā saite. Turklāt sistēmas audits atklāja, ka tie, kuriem bija uzdots sistēmu uzraudzīt, bija atbildīgi arī par tās izstrādi. Tas liecina par nepietiekamu sistēmas ārējo uzraudzību.

Melu detektori un citi rīki, kas atbalsta tiesībaizsardzības iestādes

MI sistēmas, kas atbalsta tiesībaizsardzības iestādes (piemēram, melu detektori un līdzīgi rīki), ir noteiktas kā augsta riska sistēmas Mākslīgā intelekta akta III pielikuma 6. punkta b) apakšpunktā. Viens no spilgtākajiem piemēriem, kas parāda šādu sistēmu radīto risku, ir **iBorderCtrl sistēma** automatizētu melu noteikšanas sistēma imigrācijas kontrolei.¹⁶ Sistēma tika izmantota, lai ieceļotājiem uzdotu jautājumus par viņu personīgo pieredzi un ceļojuma plāniem, tā novērtēja viņu mikroizteiksmes un ziņoja, ja radās aizdomas par negodīgumu. Šis rīks tika kritizēts kā “pseidozinātnisks”, “muļķīgs” un salīdzināts ar Orvela distopiju. Šī kritika ir balstīta uz atklājumiem, kas liecina par sistēmas neprecizitātēm un neobjektivitāti. Minētā sistēma ir aplūkota Pētījuma 5. nodaļā par MI sistēmām migrācijas, patvēruma un robežkontroles pārvaldības jomā.

VeriPol rīks nepatiesu ziņojumu par laupīšanu atklāšanai

Spānijas valsts policijas 2018. gadā ieviestā VeriPol sistēma ir vēl viens MI sistēmu, kas atbalsta tiesībaizsardzības iestādes, izmantošanas piemērs. Sistēma izmantoja MI, lai

¹⁵ [Martínez Garay, L., Boix Palop, A., Briz Redón, Á. et.al. \(2022\). Three predictive policing approaches in Spain: Viogén, RisCanvi and Veripol. Assessment from a human rights perspective.](#)

¹⁶ Skat. [Gallagher R., Jona L. \(26 July, 2019\). We tested Europe's new lie detector for travelers – and immediately triggered a false positive. The Intercept.](#)

novērtētu varbūtību, ka personas iesniegtais ziņojums par vardarbīgu laupīšanu ir nepaties. Rīks tika piemērots laupīšanas gadījumā, kas saistīta ar vardarbību vai vardarbības piedraudējumu, vai zādzības gadījumā, piemēram, kas izdarīta, uz ielas pēkšņi un spēcīgi izraujot personas rokassomiņu.¹⁷ 2024. gadā policija paziņoja, ka pārtrauc izmantot sistēmu, kā iemeslu minot, ka tās izmantošana nav noderīga tiesvedībā.¹⁸ Pētnieki norāda, ka VeriPol sistēma, visticamāk, ietvēra aizspriedumus un sniedza atšķirīgus rezultātus dažādām iedzīvotāju grupām, un sistemātiski diskriminēja dažas grupas.

Sistēma tika izveidota, pamatojoties uz klasifikāciju starp nepatiesiem un patiesiem ziņojumiem, ko sniedzis policists, tāpēc nepatiesus ziņojumus, ko šis darbinieks (vai policisti kopumā) nav spējuši atklāt, Veripol arī nevar atzīt par nepatiesiem. Turklāt Sistēma nevar atklāt, kā pilsoņi melo policijai, jo tā nedarbojas ar sūdzības iesniedzēja tekstu, ko tā sniedz policijai, bet gan ar tekstu, ko policija pieraksta no sūdzības iesniedzēja stāsta. Visbeidzot dažādie veidi, kā dažādas iedzīvotāju grupas izpauž sevi (pēc izglītības līmeņa, sociālekonomiskā statusa, ģeogrāfiskās izcelsmes, vecuma vai invaliditātes utt.), var lielā mērā ietekmēt dabiskās valodas apstrādes rīka rezultātus, un, neskatoties uz to, nav pierādījumu, ka sūdzību atlasē, ar kurām sistēma tika izveidota (ne arī vēlāk to sūdzību uzraudzībā, ar kurām tā strādā kopš tās ieviešanas policijas iecirkņos), ir bijusi kontrole pār to, vai rezultāti atšķiras atkarībā no šo dažādo grupu īpašībām. Tā rezultātā, VeriPol varētu sistemātiski diskriminēt dažas grupas. Turklāt sistēmas izmantošana netika nekādā veidā uzraudzīta. Nebija nekādas instrukcijas policijas darbiniekiem, kas skaidrotu, kad rīks ir jāizmanto vai kā būtu jārīkojas, pamatojoties uz tā sniegtajiem rezultātiem.¹⁹

Verus MI sistēma, kas atbalsta tiesībaizsardzības iestādes

Vēl viens MI sistēmas, kas atbalsta tiesībaizsardzības iestādes, piemērs ir **Verus sistēma**, ko izstrādājusi Leo Technologies.²⁰ Sistēma ir izvietota vairākos ASV cietumos, lai nodrošinātu cietumu drošību. Tā izmanto tehnoloģiju, kas pārveido runu tekstā, un kuras pamatā ir

¹⁷ [Martínez Garay, L., Boix Palop, A., Briz Redón, Á. et.al. \(2022\). Three predictive policing approaches in Spain: Viogén, RisCanvi and Veripol. Assessment from a human rights perspective.](#)

¹⁸ [García, C., Maqueda, A., Cabo D., Laursen L. \(2025\). Spanish National Police stop using Veripol, its star AI for detecting false reports.](#)

¹⁹ [Martínez Garay, L., Boix Palop, A., Briz Redón, Á. et.al. \(2022\). Three predictive policing approaches in Spain: Viogén, RisCanvi and Veripol. Assessment from a human rights perspective.](#)

²⁰ [Orakhelashvili, A. \(2024\). Real World Cases of Annex III AI Act applications that posed risks to fundamental rights.](#)

atslēgvārdi, lai uzraudzītu telefona zvanus. Tomēr šī sistēma ir radījusi bažas par pamattiesību ievērošanu. Pastāv bažas par privātumu, īpaši attiecībā uz ģimenes locekļiem vai citām personām, kas varētu sarunāties ar ieslodzītajiem, potenciāli radot jautājumus par privātuma ievērošanu. Ir paustas arī bažas par diskrimināciju, jo tehnoloģijas, kas pārveido balsi tekstā, ir neprecīzākas attiecībā uz tumšādainu personu balsīm. Turklāt, ņemot vērā cietumos esošo personu demogrāfisko sastāvu, lielāks ASV ieslodzīto īpatsvars ir tumšādainas personas un personas ar spāņu izcelsmi.²¹ Novērošanas tehnoloģiju uzraudzības projekta dibinātājs Alberts Fokss Kāns (Albert Fox Cahn) ir izteicis kritiku par šādas sistēmas izmantošanu, norādot, ka daudzi no prognozēšanas rīkiem var radīt neparedzētas kļūdas.²² Šīs tehnoloģijas tiek izmantotas, lai diskriminējoši profilētu personas cietumos un ieslodzījuma vietās, ierakstītu viņu sarunas ar advokātiem un ģimenes locekļiem un paplašinātu aizvien pieaugošo elektronisko novērošanas tīklu.²³

Kaitējuma novērtēšanas riska rīks HART atkārota noziedzīga nodarījuma iespējamības novērtēšanai

MI sistēmas, ko paredzēts izmantot tiesībaizsardzības iestādēs, lai novērtētu risku, ka fiziska persona varētu izdarīt vai atkārtoti izdarīt noziedzīgu nodarījumu, ir noteiktas kā augsta riska sistēmas saskaņā ar Mākslīgā intelekta akta III pielikuma 6. punkta d) apakšpunktu. Viens no iespējamajiem šādu MI sistēmu piemēriem ir kaitējuma novērtēšanas riska rīks HART (Harm Assessment Risk Tool – angļu val.).²⁴ HART bija algoritmisks rīks, ko Anglijā Daremas policijas iecirknis izmantoja, lai palīdzētu darbiniekiem izlemt, vai personām, kurām kā drošības līdzeklis ir piemērots apcietinājums, būtu jāpiedāvā iespēja piedalīties rehabilitācijas programmā “Checkpoint”. HART sistēma izmantoja “apcietinājuma notikumu datus”, kas iegūti no policijas apcietinājuma pārvaldības informācijas sistēmām par pieciem gadiem līdz 2012. gada beigām, un trīsdesmit četrus kritērijus riska prognozēšanai. Pamatojoties uz datiem par personas uzvedību apcietinājuma laikā un datiem no Durhamas policijas iepriekšējiem ierakstiem, HART it kā aprēķināja, kāds ir risks, ka persona izdarīs noziedzīgu nodarījumu turpmāko divu gadu laikā, ko raksturoja kā “likumpārkāpēja bīstamības”

²¹ AIAAIC. (2022). [US prison inmate AI call monitoring plan raises privacy, bias concerns.](#)

²² Mccray R. (2025). [From Surveillance to Robot Guards: How AI Could Reshape Prison Life.](#)

²³ Skat. Fassier, E. (2021). [Prison Phone Companies Are Recording Attorney-Client Calls Across the US.](#)

²⁴ Orakhelashvili, A. (2024). [Real World Cases of Annex III AI Act applications that posed risks to fundamental rights.](#)

prognozi (kas ir nepareizs apzīmējums, jo personas nebija notiesātas, un tāpēc nepareizi tika raksturotas kā “likumpārkāpējas”). Personas, par kurām tika prognozēts, ka tās, visticamāk, izdarīs “nopietnu” nodarījumu (definēts kā nodarījums, kas saistīts ar vardarbību), mazāk smagu nodarījumu vai vispār neizdarīs nodarījums, tika klasificētas attiecīgi kā “augsta”, “vidēja” vai “zema” riska personas. Tikai tās personas, kurām tika sniegta “vidēja” prognoze, bija tiesīgas izmantot Checkpoint programmu.²⁵

Viena no problēmām bija HART rīka ticamības trūkums, jo apmācības dati nebija pietiekami, lai varētu izdarīt prognozi. Tie tika veidoti no datiem par apcietinājuma kā drošības līdzekļa piemērošanu, kas pats par sevi nav noziedzības rādītājs. Līdz ar to augstāka vērtējuma piešķiršana patiesībā nebija noziedzības prognozēšana, bet labākajā gadījumā iespējamās apcietināšanas prognozēšana. Tā kā Checkpoint programmu bija tiesības izmantot tikai tām personām, kuras tika novērtētas kā “vidēja” riska, algoritmiskais lēmumu pieņemšanas process ietekmēja personu iespējas un nākotnes izredzes. Šāds rīks, kas prognozē personas turpmāko rīcību, rada tiesību uz taisnīgu tiesu un nevainīguma prezumpcijas aizskāruma risku. Turklāt, ņemot vērā šī rīka iespējamo negatīvo ietekmi uz personas tiesībām, ir svarīgi, lai personai būtu tiesības apstrīdēt procesu un lēmumus. Turklāt iespējamais iznākums, proti, personas apcietināšana, tieši ietekmē personas pamattiesības uz brīvību un drošību.

Vēl viens piemērs, kas parāda, kā MI sistēmas, kas atbalsta tiesībaizsardzības iestādes, var radīt diskriminācijas risku, ir sistēma COMPAS - Koriģējošā likumpārkāpēju pārvaldības profilēšana alternatīvām sankcijām (Correctional Offender Management Profiling for Alternative Sanction – angļu val.) potenciālā recidīvisma novērtēšanā.²⁶ Tā tiks plašāk aplūkota Pētījuma 6. nodaļā par MI sistēmām, ko izmanto tiesvedības jomā.

Prognozējošās policijas sistēmas

Mākslīgā intelekta akta III pielikuma 6. punkta d) apakšpunktā kā augsta riska MI sistēmas ir noteiktas sistēmas, kas tiek izmantotas, lai novērtētu risku, ka fiziska persona varētu izdarīt vai atkārtoti izdarīt noziedzīgu nodarījumu jeb tā sauktās prognozējošās policijas sistēmas.

²⁵ [Yeung, K., Harken, A. \(2023\). How do ‘technical’ design-choices made when building algorithmic decision-making tools for criminal justice authorities create constitutional dangers? Part II, Public Law.](#)

²⁶ [Orakhelashvili, A. \(2024\). Real World Cases of Annex III AI Act applications that posed risks to fundamental rights.](#)

Tiesībaizsardzības iestādēs MI aizspriedumi var radīt īpaši kaitīgas sekas prognozējošās policijas sistēmās. Prognozējošā policijas darbība attiecas uz analītisko metožu izmantošanu tiesībaizsardzības iestādēs, lai veiktu statistiskas prognozes par iespējamu noziedzīgu darbību. Tā ir policijas stratēģija, kas izmanto algoritmisku uzraudzību, lai prognozētu turpmākus noziedzīgus nodarījumus, noziedzniekus un upurus, lai iejauktos pirms noziedzīgu nodarījumu izdarīšanas.²⁷ Var izšķirt divus prognozējošās policijas darbības veidus: paredzošo kartēšanu un paredzošo identifikāciju. Visbiežāk izmantotais prognozējošās policijas darbības veids ir paredzošā kartēšana vai uz vietu balstīta paredzošā policijas darbība. Tas attiecas uz progresīvu ģeotelpisko analīzi, lai prognozētu, kad un kur noziedzīgs nodarījums varētu notikt apkopotā analīzes līmenī. Apkopotā analīzes līmenī tiek veikta noziedzības datu kolektīva izpēte noteiktos ģeogrāfiskos apgabalos vai laika periodos, lai identificētu modeļus un tendences. Šī pieeja balstās uz noziedzīga nodarījuma vietas vai paaugstināta riska laika prognozēšanu, nevis uz konkrētu personu vai notikumu mērķēšanu. Ja uz vietu balstīta paredzošā policijas darbība, reaģējot uz augstāku noziedzīgu nodarījumu biežumu dažās apkaimēs, vairāk mērķē uz grupām, kas saistītas ar aizliegtajiem diskriminācijas kritērijiem, tā ir netieša diskriminācija. Piemēram, paredzošās policijas darbības algoritmu izmantošana var izraisīt nevajadzīgi palielinātu policijas klātbūtni vietās, kurās galvenokārt dzīvo noteiktas etniskās minoritātes. Savukārt paredzošā identifikācija ir analīze personas vai grupas līmenī, un šādā gadījumā tiek apstrādāti personas dati. Tā var koncentrēties uz potenciālo likumpārkāpēju, likumpārkāpēju identitātes, kriminālās uzvedības un potenciālo noziedzīga nodarījuma upuru prognozēšanu. Kā tika norādīts iepriekš, Mākslīgā intelekta akts ierobežo prognozējošās policijas darbības. Lai mazinātu riskus, tai skaitā nepamatotu arestu iespējamību, būtiska nozīme ir cilvēka veiktai uzraudzībai. Standarta prakse ir jebkuras MI sakritības pārbaude, ko veic vismaz divi cilvēki.

Eiropā ir vairāki piemēri, kad tiek izmantotas prognozējošās policijas sistēmas, kas iespējams rada diskriminācijas risku. Viens no piemēriem ir Nīderlandes Sensing projekts, kas tika veikts no 2019. gada janvāra līdz 2020. gada oktobrim, un bija vērsts uz tādu noziedzīgu nodarījumu kā veikalu zādzības apkarošanu dienvidaustrumu pilsētā Rūrmondā. Sensing projekts izmantoja attālinātos sensorus pilsētā un tās apkārtnē, lai noteiktu to automašīnu

²⁷ [Europol \(2025\). AI bias in law enforcement - A practical guide, Europol Innovation Lab observatory report.](#)

marku, krāsu un maršrutu, kurās pārvadāja personas, kuras tika turētas aizdomās par to, ko policija dēvē par “mobilo bandītismu”. Sensing projekts identificēja transportlīdzekļus ar Austrumeiropas numura zīmēm, mēģinot atšķirt romus kā iespējamus kabatzagļus un veikalu zagļus. Pats modelis bija īpaši vērsts pret personām, kas nav Nīderlandes pilsoņi. Rūrmondas policija izslēdza Nīderlandes pilsoņus no “mobilā bandītisma” definīcijas un sašaurināja Sensing projekta jomu, kas bija nepārprotami neobjektīvi un diskriminējoši pret noteiktām iedzīvotāju grupām uz rases un etniskās piederības pamata. Konkrētāk, sistēma bija vērsta uz augstāku riska rādītāju noteikšanu personām ar Austrumeiropas un romu etnisko piederību, kā rezultātā šīm grupām bija lielāka iespējamība tikt pakļautai tādiem pasākumiem kā datu glabāšana policijas datubāzēs.²⁸

Vēl viens piemērs ir Vācijā un Šveicē izmantotā prognozējošā kartēšanas programmatūra PRECOBS.²⁹ Tā izmanto algoritmus un informāciju par pagātnē izdarītiem noziedzīgiem nodarījumiem, lai prognozētu tā saukto “gandrīz atkārtotu” noziedzīgu nodarījumu izdarīšanu. Izmantojot nesenus pagātnes pārkāpumu datus, policijas iestādēm tiek sniegtas prognozes par noteiktu “rajonu”, un tās tiek izmantotas operatīviem pasākumiem un noziedzīgu nodarījumu novēršanai, tādējādi radot iespējamu diskriminācijas risku pret noteiktām iedzīvotāju grupām, kas galvenokārt tur dzīvo. Prognozējošā policijas darbība ir paredzēta, lai sniegtu prognozes par konkrētiem noziedzīgiem nodarījumiem (piemēram, zādzību, laupīšanu, ļaunprātīgu dedzināšanu). Prognozēšanas programmatūru izstrādāja Modeļu prognozēšanas tehnoloģiju institūts (The Institute for Pattern-based Forecasting Technology IfmPt), kuru 2021. gadā pārņēma Logobject Deutschland GmbH. Vācijas policija Karlsrūē un Štutgartē nolēma pārtraukt PRECOBS programmatūras izmantošanu, jo nebija pietiekamu noziedzīgu nodarījumu datu, lai veiktu ticamas prognozes.

Viens no galvenajiem izaicinājumiem prognozējošajā policijas darbā ir neobjektīvu datu izmantošana. Algoritmu apmācībai izmantotie dati ir vēsturiski dati no policijas datubāzēm, un tie var nebūt reprezentatīvi pašreizējā brīdī. Apkopotie dati ir divējādi: no ziņotajiem noziedzīgajiem nodarījumiem un no noziedzīgiem nodarījumiem, ko novērojusi vai atklājusi

²⁸ [Amnesty International. \(2020\). We sense trouble: Automated Discrimination and Mass Surveillance in Predictive Policing in the Netherlands.](#)

²⁹ [Egbert, S., Krasmann, S. \(2019\). Predictive policing: Not yet, but soon preemptive? *Policing and Society*, 30\(8\), 905–919.](#)

pati policija. Datu izplatību ietekmē noteiktu sociālekonomisko grupu nevēlēšanās ziņot par noziegumiem, palielināta policijas klātbūtne noteiktās teritorijās, noziedzīga nodarījuma veids, kas padara to vairāk vai mazāk “novērojamu” vai, visticamāk, par ko tiks ziņots utt. Turklāt var pastāvēt neobjektivitāte apmācību procesā, ja pats mašīnmācīšanās modelis var pastiprināt jau esošo neobjektivitāti. Visbeidzot, izlases neobjektivitāte rodas tāpēc, ka vairāk noziedzīgie nodarījumi tiek novēroti tur, kur vairāk atrodas policija. Pastāvot šīm neobjektivitātēm un tās nesamazinot, pēc noteikta laika neobjektīvās MI sistēmas izejas dati tiek atgriezti kā ievades dati atgriezeniskās saites cilpas mehānisma dēļ. Sistēmas izmantošanas rezultātā lielākā daļa policijas spēku tiek norīkoti vienā apgabalā, kurā vēsturiski ir bijis vairāk noziedzīgu nodarījumu nekā citā. Tas noved pie tādu datu kopu izveides, kas šķietami atspoguļo augstāku noziedzības līmeni, bet patiesībā atspoguļo lielāku policijas uzmanību. Turklāt paši mašīnmācīšanās modeļi var veidot atgriezeniskās saites cilpas, ja tie turpina mācīties pēc ieviešanas. Tipiski piemēri ir apmācības modeļi, kas periodiski tiek apmācīti, izmantojot datus, kas uzkrāti grupas vai tiešsaistes mācību modeļos, kas tiek nepārtraukti apmācīti, saņemot jaunus datus.

Šāds algoritms, piemēram, tika izmantots Nīderlandes bērnu aprūpes pabalstu krāpšanas skandālā, kas ir viens no spilgtākajiem piemēriem, kas atklāj, kā MI sistēmu izmantošana var radīt diskrimināciju. Nīderlandes valdība ieviesa sistēmiskā riska identificēšanas (SyRI) krāpšanas apkarošanas sistēmu, ko izmantoja, lai atklātu krāpšanu sociālās labklājības jomā. Tās izmantošanas rezultātā vairāk kā divdesmit tūkstoši vecāku tika nepamatoti apsūdzēti par krāpnieciskiem pabalstu pieprasījumiem, pieprasot viņiem atmaksāt saņemtos pabalstus. Daudzos gadījumos summa sasniedza vairākus desmitus tūkstošu eiro, radot ģimenēm smagas finansiālās grūtības. Turklāt SyRI sistēmu, galvenokārt, izmantoja apkaimēs, kurās dzīvoja maznodrošinātas personas, un tā radīja diskrimināciju iedzīvotāju sociālekonomiskā stāvokļa un etniskās izcelsmes dēļ. 2020. gadā Hāgas tiesa lika apturēt SyRI darbību, secinot, ka tiesību akti, ar kuriem tika izveidots SyRI, nenodrošināja pietiekamu aizsardzību pret iejaukšanos privātajā dzīvē, jo tika veikti nesamērīgi pasākumi krāpšanas novēršanai un sodīšanai ekonomiskās labklājības interesēs.³⁰ Syri sistēma un vairāki citi

³⁰ [UNESCO. \(2023\). Global toolkit on AI and the rule of law for the judiciary.](#); [Algorithm Watch. \(2020\). How Dutch activists got an invasive fraud detection algorithm banned.](#); [Murad M. \(2022\). Fighting Back on Algorithmic Opacity.](#)

prakses piemēri, kas parāda riskus, ko rada sistēmas, ko iestādes izmanto krāpšanas atklāšanai sociālās labklājības jomā, ir aplūkoti 2024. gada Pētījumā.

4. Kritiskā infrastruktūra

Saskaņā ar Mākslīgā intelekta aktu par augstu riska MI sistēmām uzskata sistēmas, ko izmanto kritiskās infrastruktūras jomā – MI sistēmas, ko paredzēts izmantot kā drošības sastāvdaļas kritiskās digitālās infrastruktūras pārvaldībā un darbībā, ceļu satiksmē vai ūdens, gāzes vai elektroapgādē vai apkures nodrošināšanā (6. panta 2. punkts, III pielikuma 2. punkts).

Mākslīgā intelekta akts skaidro, ka nenovēršamus draudus fizisku personu dzīvībai vai fiziskajai drošībai var radīt Eiropas Parlamenta un Padomes Direktīvas (ES) 2022/2557³¹ 2.panta 4. punktā definētās kritiskās infrastruktūras nopietni traucējumi, ja šādas kritiskās infrastruktūras traucējumu vai iznīcināšanas rezultātā nenovēršami tiktu apdraudēta personas dzīvība vai fiziskā drošība, tostarp, nopietni kaitējot pamatapgādes nodrošināšanai iedzīvotājiem vai valsts pamatfunkciju īstenošanai (33. apsvērums).

Attiecībā uz kritiskās infrastruktūras pārvaldību un darbību kā augsta riska sistēmas tiek klasificētas tās MI sistēmas, kuras paredzēts lietot kā drošības sastāvdaļas Direktīvas (ES) 2022/2557 pielikuma 8. punktā uzskaitītās kritiskās digitālās infrastruktūras pārvaldībā un darbībā, ceļu satiksmē, kā arī ūdens, gāzes un elektroenerģijas piegādē un apkures nodrošināšanā, jo sistēmu darbības atteice vai traucējumi var plašā mērogā apdraudēt cilvēku veselību un dzīvību un izraisīt ievērojamus traucējumus sociālo un saimniecisko darbību parastajā norisē. Kritiskās infrastruktūras, tostarp kritiskās digitālās infrastruktūras, drošības sastāvdaļas ir sistēmas, ko lieto, lai tieši aizsargātu kritiskās infrastruktūras fizisko integritāti vai cilvēku veselību un drošību un mantas drošību, bet kas nav nepieciešamas sistēmas darbībai. Šādu sastāvdaļu atteice vai darbības traucējumi var tieši radīt riskus kritiskās infrastruktūras fiziskajai integritātei un tādējādi riskus cilvēku veselībai un drošībai un mantas drošībai (Mākslīgā intelekta akta 55. apsvērums).

³¹ [Eiropas Parlamenta un Padomes Direktīva \(ES\) 2022/2557 \(2022. gada 14. decembris\) par kritisko vienību noturību un Padomes Direktīvas 2008/114/EK atcelšanu. OV L 333, 27.12.2022.](#)

Mākslīgā intelekta akta preambulā arī skaidrots, ka sastāvdaļas, ko paredzēts lietot tikai kibernetikas nolūkos, nav uzskatāmas par drošības sastāvdaļām. Šādas kritiskās infrastruktūras drošības sastāvdaļu piemēri var būt sistēmas ūdens spiediena uzraudzībai vai ugunsgrēka signalizācijas sistēmas mākoņdatošanas centros (55. apsvēruma).

ES kritiskās infrastruktūras tiesību akti (Direktīva (ES) 2022/2557³² un Regula (ES) 2023/2450³³) nosaka plašu nozaru sarakstu, kas ietver komunālos pakalpojumus un enerģētiku, ūdensapgādi, transportu, pārtiku, atkritumu apsaimniekošanu, sabiedrisko pakalpojumu pārvaldi, komunikāciju digitālo infrastruktūru un kosmosu.

MI sistēmu izmantošana kritiskajā infrastruktūrā tādējādi ir saistīta ar Mākslīgā intelekta akta III pielikuma 5. punktu, kas kā augsta riska MI jomu nosaka piekļuvi privātiem pamatpakalpojumiem un sabiedriskajiem pamatpakalpojumiem un pabalstiem un to izmantošanu.³⁴ 2024. gada Pētījumā tika detalizētāk aplūkota minētā joma, kā arī analizēti prakses piemēri, kas atklāj diskriminācijas riskus.

MI sistēmas, ko izmanto kritiskās infrastruktūras jomā, var apdraudēt personu veselību un drošību un nav pamatā saistīti ar diskriminācijas risku. Lai gan konkrēti prakses piemēri, kas atklātu, kā var pārkāpt diskriminācijas aizlieguma principu, nav atrodam, tajā pašā laikā MI izmantošana šajā jomā rada bažas par kibernetikas riskiem MI vadītā kritiskās infrastruktūras aizsardzībā. Praksē pastāv arī citi ētiski un tiesiski izaicinājumi, piemēram, kas saistīti ar MI sistēmas izmantošanu kritiskās infrastruktūras pamatpakalpojumu pārvaldību un rada bažas par diskriminējošu resursu sadali un cenu noteikšanu. Zemāk ir sniegti dažādi piemēri, kas atklāj šos riskus.

³² [Eiropas Parlamenta un Padomes Direktīva \(ES\) 2022/2557 \(2022. gada 14. decembris\) par kritisko vienību noturību un Padomes Direktīvas 2008/114/EK atcelšanu. OV L 333, 27.12.2022.](#)

³³ [Komisijas deleģētā regula \(ES\) 2023/2450 \(2023. gada 25. jūlijs\), ar ko papildina Eiropas Parlamenta un Padomes Direktīvu \(ES\) 2022/2557, izveidojot pamatpakalpojumu sarakstu/. OV L, 2023/2450, 30.10.2023.](#)

³⁴ [Mazzarella J. \(2024\). AI in Critical Infrastructure Markets: Are Smart Systems AI? The EU Act Says They May Be. AI Alert.](#)

Kiberdrošības un ētikas izaicinājumi saistībā ar kritiskās informācijas aizsardzību MI balstītā vidē

MI integrācija paplašina uzbrukuma iespējas kritiskajā enerģētikas infrastruktūrā, padarot sistēmas par pievilcīgiem mērķiem ar potenciāli postošām sekām.³⁵ Izaicinājumi ietver ar MI saistītas ievainojamības, piemēram, pretinieku uzbrukumus un datu saindēšanu³⁶, riskus, kas saistīti ar MI integrēšanu ar novecojušām sistēmām un lietu interneta ierīcēm³⁷, potenciālus datu privātuma pārkāpumus, kas grauj sabiedrības uzticēšanos³⁸, un draudus sistēmas stabilitātei un drošībai³⁹. Pat MI vadītiem kiberdrošības rīkiem ir ievainojamības, un drošības draudus palielina darbinieku prasmju trūkums.

MI izmantošana, lai pieņemtu lēmumus, kas saistīti ar enerģētiku, rada ētiskas bažas par algoritmisko neobjektivitāti, datu privātumu un atbildību.⁴⁰ MI sistēmas, kas apmācītas, izmantojot vēsturiskus datus, kas atspoguļo sabiedrības nevienlīdzību⁴¹, var novest pie diskriminējošiem rezultātiem resursu sadalē vai cenu noteikšanā, graujot sabiedrības uzticēšanos⁴². Datu privātums, izmantojot viedos skaitītājus un lietu internetu, ir svarīgs, jo pārkāpumi ir saistīti ar kiberdrošības riskiem. Savukārt MI izskaidrojamības trūkums sarežģī jautājumu par atbildību par MI sistēmu kļūdām.

Viens no MI sistēmu ieviešanas gadījumiem, kas ir radījis bažas, ir saistīts ar būtisku resursu, piemēram, ūdens un elektroenerģijas piegādi, pat tad, ja sistēmas ir ieviestas sociāli atbalstāmiem mērķiem. Spānijas valdība ieviesa programmu **Bono Social de Electricidad**

³⁵ Park, C. (2025). Addressing Challenges for the Effective Adoption of Artificial Intelligence in the Energy Sector. *Sustainability*, 17(13), 5764.

³⁶ Skat. Mengidis, N., Tsikrika, T., Vrochidis, S. et.al. (2019). Blockchain and AI for the next generation energy grids: Cyber-security challenges and opportunities. *Information & Security*, 43 (1), p. 21–33.

³⁷ Strielkowski, W., Vlasov, A., Selivanov K, et.al. (2023). Prospects and Challenges of the Machine Learning and Data-Driven Methods for the Predictive Analysis of Power Systems: A Review. *Energies*. 16(10), 4025.

³⁸ Kumari, A., Gupta, R., Tanwar, S. et.al. (2020). Blockchain and AI amalgamation for energy cloud management: Challenges, solutions, and future directions. *J. Parallel Distrib. Comput.* 143, p. 148–166.

³⁹ Shi, Z.; Yao, W.; Li, Z.; Zeng, L. et.al. (2020). Artificial intelligence techniques for stability analysis and control in smart grids: Methodologies, applications, challenges and future directions. *Appl. Energy*, 278.

⁴⁰ Park, C. (2025). Addressing Challenges for the Effective Adoption of Artificial Intelligence in the Energy Sector. *Sustainability*, 17(13), 5764.

⁴¹ Nguyen, V.N.; Tarefko, W.; Sharma, P. et.al. (2024). Potential of explainable artificial intelligence in advancing renewable energy: Challenges and prospects. *Energy Fuels*, 38, p. 1692–1712.

⁴² Turpat.

(BSE), kuras mērķis bija **nodrošināt atlaižu piešķiršanu par enerģijas rēķiniem riska grupas personām un ģimenēm** Spānijā, lai mazinātu enerģētisko nabadzību. Šādai MI sistēmai, kas pārvalda atlaižu tiesību piešķiršanu, ir potenciāls MI diskriminācijas risks, jo atkarībā no apmācības datiem, tās var ietvert neobjektivitāti, radot netaisnīgus lēmumus neaizsargātām iedzīvotāju grupām, atspoguļojot pastāvošo nevienlīdzību vai radot jaunu. Sistēmas novērtējums liecināja, ka maksājumu atlaides nemazināja enerģētisko nabadzību un, iespējams, to ir palielinājušas. Turklāt sistēma tika kritizēta par nepārredzamību.⁴³ Rūpīga ētiska MI pārvaldība ir būtiska, lai nodrošinātu taisnīgu, atbildīgu un pārredzamu algoritmisko lēmumu pieņemšanu.

Attiecībā uz kritiskās informācijas aizsardzību pret kiberdrošības riskiem, padziļinātas aizsardzības stratēģijas ir svarīgas, taču tās rada organizatoriskus un ekonomiskus izaicinājumus. Kritiskās infrastruktūras īpašniekiem un valdītājiem vispirms ir jāsaprot, kā, kur un kāpēc MI sistēmas tiks izmantotas, lai novērtētu kontekstam un nozarei specifiskus riskus un iespējamo ietekmi uz drošību un aizsardzību. Tas ietver pienākumu dokumentēt potenciālo kontekstam specifisko ietekmes uz drošību un aizsardzību novērtējumu, kas saistīts ar MI sistēmām un kas ietekmē personas, kopienas un sabiedrību, iekļaujot potenciālos riskus, kas saistīti ar aizspriedumiem, privātumu un sensitīvas informācijas ļaunprātīgu izmantošanu.⁴⁴

⁴³ [Orakhelashvili, A. \(2024\). Real World Cases of Annex III AI Act applications that posed risks to fundamental rights.](#); [García Alvarez, G., Tol, R. \(2023\). The impact of the Bono Social de Electricidad on energy poverty in Spain. University of Sussex.](#)

⁴⁴ [US Department of Homeland Security. \(2024\). Mitigating Artificial Intelligence \(AI\) Risk: Safety and Security Guidelines for Critical Infrastructure Owners and Operators.](#)

5. Migrācijas, patvēruma un robežkontroles pārvaldība

Vēl viena joma, kurā MI sistēmu izmantošana rada augstu risku, ir migrācijas, patvēruma un robežkontroles pārvaldība, ciktāl to izmantošanu atļauj attiecīgie ES vai valsts tiesību akti (Mākslīgā intelekta akta 6. panta 2. punkts, III pielikuma 7. punkts). Augstu risku rada MI sistēmas migrācijas, patvēruma un robežkontroles pārvaldības jomā, ko kompetentās publiskajās iestādēs vai to vārdā, vai ES iestādēs, struktūrās, birojos vai aģentūrās paredzēts izmantot šādos četros veidos:

- 1) kā melu detektorus vai līdzīgus rīkus;
- 2) lai novērtētu risku, tostarp drošības risku, neatbilstīgas migrācijas risku vai veselības risku, ko rada fiziska persona, kas vēlas ieceļot vai ir ieceļojusi dalībvalsts teritorijā;
- 3) lai palīdzētu kompetentām publiskajām iestādēm izskatīt patvēruma, vīzu vai uzturēšanās atļauju pieteikumus, kā arī saistītās sūdzības, nosakot, vai fiziskās personas, kuras iesniegušas pieteikumu uz kādu no minētajiem statusiem, atbilst kritērijiem, tostarp veicot saistītu pierādījumu patiesuma novērtēšanu;
- 4) saistībā ar migrācijas, patvēruma vai robežkontroles pārvaldību nolūkā atklāt, atpazīt vai identificēt fiziskas personas, izņemot ceļošanas dokumentu verifikāciju.

Mākslīgā intelekta akts vērš uzmanību uz to, ka migrācijas un patvēruma jomā un robežkontroles pārvaldībā lietotās MI sistēmas ietekmē personas, kuras bieži ir īpaši neaizsargātā pozīcijā un ir atkarīgas no kompetento publisko iestāžu darbības rezultāta. Tāpēc šajā jomā lietoto MI sistēmu precizitāte, nediskriminējošais raksturs un pārredzamība ir īpaši svarīga, lai tiktu ievērotas personu pamattiesības, jo īpaši tiesības brīvi pārvietoties, tiesības uz nediskriminēšanu, privātums un personas datu aizsardzība, tiesības uz starptautisko aizsardzību un labu pārvaldību (60. apsvērums).

Mākslīgā intelekta akts skaidro, ka MI sistēmām migrācijas, patvēruma un robežkontroles pārvaldības jomā, uz kurām attiecas Mākslīgā intelekta akts, ir jāatbilst attiecīgajām procesuālajām prasībām, kas noteiktas Eiropas Parlamenta un Padomes Regulā (EK)

Nr. 810/2009 (Vīzu kodekss)⁴⁵, Eiropas Parlamenta un Padomes Direktīvā 2013/32/ES⁴⁶ un citos ES tiesību aktos. Dalībvalstīm vai ES iestādēm, struktūrām, birojiem vai aģentūrām nekādā gadījumā nebūtu jālieto MI sistēmas migrācijas, patvēruma un robežkontroles pārvaldībā kā līdzekli, lai apietu savas starptautiskās saistības saskaņā ar 1951. gada 28. jūlijā Ženēvā parakstīto ANO Konvenciju par bēgļa statusu, kas grozīta ar 1967. gada 31. janvāra protokolu. Tās arī nebūtu jālieto, lai jebkādā veidā pārkāptu neizraidīšanas principu vai liegtu drošas un efektīvas likumīgas iespējas iekļūt ES teritorijā, tostarp tiesības uz starptautisko aizsardzību (60. apsvērums).

Pētījuma 3. nodaļā, aplūkojot augsta riska sistēmas tiesībaizsardzības jomā, jau tika vērsta uzmanība uz nepieciešamību nošķirt augsta riska MI sistēmas no aizliegtas MI prakses. Šāda MI sistēmu nošķiršana ir būtiska arī migrācijas, patvēruma un robežkontroles pārvaldības jomā.

Saskaņā ar Mākslīgā intelekta akta 5. panta 1. punkta d) apakšpunktu ir aizliegta MI sistēmu lietošana fizisku personu radītā riska novērtējuma veikšanai, kuru mērķis ir novērtēt vai prognozēt personas noziedzīga nodarījuma izdarīšanas risku, balstoties vienīgi uz personas profilēšanu (Mākslīgā intelekta akta 3. panta 52. punkts un VDAR⁴⁷ 4. panta 4. punkts) vai uz personas īpašību un personības novērtējumu.

Šis aizliegums atbilst nevainīguma prezumpcijas principam un ir būtisks Vīzu kodeksa kontekstā, kur trešās valsts valstspiederīgajam var atteikt ieceļošanu vai vīzas piešķiršanu, ja šī persona tiek uzskatīta par draudu sabiedriskajai kārtībai vai iekšējai drošībai (Vīzu kodeksa 32. panta 1. punkta a) apakšpunkta vi) punkts). Lai gan profilēšana ir atļauta izņēmuma kārtā saskaņā ar VDAR, MI izmantošana šim nolūkam ir aizliegta. Personu var atzīt par aizdomās turēto noziedzīga nodarījuma izdarīšanā tikai tad, ja šādas aizdomas ir balstītas uz cilvēka veiktu objektīvu faktu novērtējumu.⁴⁸ Mākslīgā intelekta akta 42. apsvērums skaidro, ka

⁴⁵ [Eiropas Parlamenta un Padomes Regula \(EK\) Nr. 810/2009 \(2009. gada 13. jūlijs\), ar ko izveido Kopienas Vīzu kodeksu \(Vīzu kodekss\), OV L 243, 15.9.2009.](#)

⁴⁶ [Eiropas Parlamenta un Padomes Direktīva 2013/32/ES \(2013. gada 26. jūnijs \) par kopējām procedūrām starptautiskās aizsardzības statusa piešķiršanai un atņemšanai \(pārstrādāta redakcija\). OV L 180, 29.6.2013.](#)

⁴⁷ [Eiropas Parlamenta un Padomes Regula \(ES\) 2016/679 \(2016. gada 27. aprīlis\) par fizisku personu aizsardzību attiecībā uz personas datu apstrādi un šādu datu brīvu apriti un ar ko atceļ Direktīvu 95/46/EK \(Vispārīgā datu aizsardzības regula\).](#)

⁴⁸ [Pina M.C. \(2025\). AI in the context of border management, migration and asylum in the EU: technological innovation vs. fundamental rights of migrants in the AI Act.](#)

saskaņā ar nevainīguma prezumpciju fiziskas personas ES vienmēr būtu jāvērtē pēc to faktiskās uzvedības. Fiziskas personas nekad nevajadzētu vērtēt pēc MI prognozētas uzvedības, pamatojoties tikai uz to profilēšanu, personiskajām īpašībām vai rakstura iezīmēm, piemēram, valstspiederību, dzimšanas vietu, dzīvesvietu, bērnu skaitu, parāda līmeni vai automobiļa veidu, ja, balstoties uz objektīviem pārbaudāmiem faktiem, nav pamatotu aizdomu par personas iesaistīšanos noziedzīgā darbībā un ja nav cilvēka veikta novērtējuma par to. Tāpēc būtu jāaizliedz riska novērtējumi, ko veic attiecībā uz fiziskām personām, lai novērtētu to nodarījuma izdarīšanas risku vai lai prognozētu noziedzīga nodarījuma notikšanu, pamatojoties tikai uz personu profilēšanu vai novērtējot to personiskās īpašības un rakstura iezīmes.

Mākslīgā intelekta akts aizliedz arī MI sistēmu izmantošanu sejas atpazīšanas datubāžu izveidei vai paplašināšanai, nejauši apkopojot attēlus no interneta vai videonovērošanas kameru ierakstiem (5. panta 1. punkta e) apakšpunkts), lai aizsargātu privātumu un novērstu masveida novērošanu. Mākslīgā intelekta akts aizliedz arī fizisku personu biometriskās kategorizācijas sistēmas, kuru pamatā ir viņu biometriskie dati, lai atvedinātu vai izsecinātu viņu rasi, politiskos uzskatus, dalību arodbiedrībās, reliģisko vai filozofisko pārliecību, seksuālo dzīvi vai orientāciju (5. panta 1. punkta g) apakšpunkts), izņemot likumīgi iegūtu biometrisku datu kopu marķēšanu un filtrēšanu tiesībaizsardzības jomā. Šī atšķirība ir svarīga, jo ES dalībvalstis arvien vairāk izmanto tehnoloģijas, lai pārbaudītu patvēruma meklētāju drošību un identitāti.⁴⁹ Mākslīgā intelekta akts aizliedz arī reāllaika attālinātās biometriskās identifikācijas sistēmas publiski pieejamās vietās, izņemot gadījumus, kad tas tiek uzskatīts par nepieciešamu konkrētiem mērķiem, kas izklāstīti Mākslīgā intelekta aktā (5. panta 1. punkta h) apakšpunkts, 2. punkts).

Jāatzīmē, ka aizliegto MI lietojumu un sistēmu saraksts nav izsmeļošs, un šīs prakses var aizliegt arī citi tiesību akti. Vispārējie aizliegumi, piemēram, diskriminācijas aizliegums, attiecas uz Mākslīgā intelekta aktu.

Tajā pašā laikā praksē pastāvīgais ES informācijas sistēmu atjaunināšanas, paplašināšanas un savstarpējas savienošanas process robežu pārvaldībai ir radījis jaunas iespējas testēt un

⁴⁹ Turpat.

ievieš MI tehnoloģijas uz robežām. Piemēram, lielākā daļa ES informācijas sistēmu ir paplašinātas, lai varētu apkopot sejas attēlus, kas ļautu izmantot sejas atpazīšanas tehnoloģijas. Šādas tehnoloģijas arvien biežāk tiek izmantotas lidostās, lai automātiski pārbaudītu ceļotāju identitāti. Jaunā Eiropas ceļošanas informācijas un atļauju sistēma (ETIAS), kas pašlaik tiek izstrādāta, varētu ietvert papildu automatizācijas vai analītikas līmeni, kas balstīts uz MI vai mašīnmācīšanos, lai identificētu aizdomīgus pieteikumus.⁵⁰ Šo sistēmu atbilstības nodrošināšana Mākslīgā intelekta akta prasībām ir īpaši svarīga.

Augsta riska MI sistēmas, lai gan nav aizliegtas, var radīt nopietnus riskus personu veselībai, drošībai vai pamattiesībām (Mākslīgā intelekta akta 48. apsvērums). Ņemot vērā šo MI sistēmu iespējamās nopietnās sekas, tām ir jāatbilst stingrākām prasībām, kas ir noteiktas Mākslīgā intelekta aktā, piemēram, riska pārvaldība (9. pants), pārredzamība (13. pants), cilvēka virsvadība (14. pants), kiberdrošība, precizitāte un noturība (15. pants), datu kvalitāte un pārvaldība, kā arī sistēmu apmācība un testēšana (10. pants), reģistrācija ES datubāzē (49. un turpmākie panti) un iepriekšējs ietekmes uz pamattiesībām novērtējums (27. pants).

Mākslīgā intelekta akts papildus uzliktajiem pienākumiem aizsargā to cilvēku tiesības, kurus skar šādas sistēmas. Jo īpaši tas garantē tiesības uz skaidriem un jēgpilniem skaidrojumiem par MI sistēmas lomu lēmumu pieņemšanas procesā un galvenajiem lēmuma elementiem (86. pants). Tas kalpo kā garantija tiesībām uz efektīvu tiesisko aizsardzību (Hartas 47. pants). To ir apstiprinājusi Eiropas Savienības Tiesa (EST) *Ligue des Droits Humains* lietā, kurā tika aplūkota automatizēta riska novērtēšana, kuras pamatā ir Pasažieru datu reģistra sistēma. EST uzsvēra nepieciešamību pēc pārredzamības un piekļuves informācijai, nosakot, ka personai ir jāspēj saprast, kā darbojas lēmumu kritēriji un izmantotā programma, lai, pilnībā zinot attiecīgos faktus, varētu izlemt, vai apstrīdēt to nelikumīgo vai diskriminējošo raksturu.⁵¹

Pētnieki norāda, ka Mākslīgā intelekta akts ir nepilnīgs, kas var radīt potenciāli negatīva ietekmi uz migrantu pamattiesībām. Gan aizliegto MI sistēmu saraksts, gan augsta riska

⁵⁰ [European Parliament. \(2025\). Artificial intelligence in asylum procedures in the EU.](#)

⁵¹ Eiropas Savienības Tiesas 2022. gada 21. jūnija spriedums lietā C-817/19 *Ligue des droits humains*, CLI:EU:C:2022:491

sistēmu saraksts ir nepilnīgs, un pastāv problēmas ar pārredzamības un cilvēka uzraudzības nodrošināšanu migrācijas jomā. Tiek ieteikts veikt stingrākus pasākumus, piemēram, aizliegt aizskarošas un zinātniski nepamatotas sistēmas, paplašināt augsta riska kategorijas, kā arī nodrošināt cilvēka uzraudzību un uz to pamata pieņemto lēmumu pārredzamību. Mākslīgā intelekta aktā ir paredzēta iespēja Eiropas Komisijai pārskatīt aizliegto MI sistēmu sarakstu, kā arī Mākslīgā intelekta akta III pielikumu, lai atjauninātu augsta riska MI sistēmu sarakstu atbilstoši tehnoloģiju attīstībai un sociālajām vajadzībām (7., 97. un 112. pants).

Turpmāk ir minēti vairāki uzskatāmi piemēri no ārvalstu prakses, kas parāda, kā MI sistēmu izmantošana migrācijas, patvēruma un robežkontroles pārvaldības jomā var pārkāpt diskriminācijas aizlieguma principu.

Prakses piemēri

iBorderCtrl sistēma

Saskaņā ar Mākslīgā intelekta aktu migrācijas, patvēruma un robežkontroles pārvaldības jomā, augstu risku rada MI sistēmas, ko paredzēts izmantot publiskajās iestādēs kā melu detektorus un līdzīgus rīkus (III pielikuma 7. punkta a) apakšpunkts). Līdzīgi arī tiesībaizsardzības jomā augstu risku rada MI sistēmas, kas atbalsta tiesībaizsardzības iestādes, piemēram, melu detektori un līdzīgi rīki (Mākslīgā intelekta akta III pielikuma 6. punkta b) apakšpunkts).

Spilgts piemērs, kas parāda šādu sistēmu radīto risku, ir iBorderCtrl sistēma. ES finansētais projekts iBorderCtrl, kurā tika paredzēts izstrādāt automatizētu melu noteikšanas sistēmu imigrācijas kontrolei, izraisīja plašu kritiku un tika pārtraukts 2019. gadā.⁵² Projektā iesaistījās Ungārijas, Grieķijas, kā arī Latvijas drošības iestādes. Projekts paredzēja ieviest ieceļotāju kontroles sistēmu, kas darbotos šādi – ieceļotājs, kas vēlas iekļūt ES, pirms ierašanās lidostā, izmantojot savu datoru, piesakās vietnē, augšupielādējot savas pasas attēlu. Uz ekrāna parādās virtuāls policists tumši zilā formas tērpā. Viņš uzdod dažādus jautājumus, piemēram: “Kāds ir tavs uzvārds?”, “Kāda ir tava pilsonība”, “Kāds ir tavs

⁵² [Stolton, S. \(8 February, 2021\). Commission under pressure in EU court over ‘lie detector tech’. Euractiv.](#); Skat. arī [Gallagher R., Jona L. \(26 July, 2019\). We tested Europe’s new lie detector for travelers – and immediately triggered a false positive. The Intercept.](#)

ceļojuma mērķis?”. Jūs atbildat uz uzdotajiem jautājumiem, un virtuālais policists izmanto jūsu tīmekļa kameru, lai skenētu seju un acu kustības, lai noteiktu, vai melojat vai nē.

Intervijas beigās sistēma piešķir kvadrātkodu, kas jāuzrāda apsargam robežkontrolē. Apsargs noskenē kodu, izmantojot planšetdatoru, noņem pirkstu nospiedumus un pārskata uzņemto sejas attēlu, lai pārbaudītu, vai tas atbilst pasei. Apsarga planšetdatorā tiek parādīts rezultāts 100 punktu skalā, kas parāda, vai mašīna ir atzinusi sacīto par patiesību.

Ceļotājiem, kurus uzskata par bīstamiem, var liegt ieceļošanu ES. The Intercept veiktais pētījums atklāja, ka 4 no 16 godīgām atbildēm tika nepareizi identificētas kā meli.⁵³

2019. gadā Eiropas Savienības Tiesas Vispārējā tiesā iesniegtā Eiropas Parlamenta deputāta Patrika Breijera prasība (lieta T-158/19) pret Eiropas Komisiju⁵⁴ rada būtiskus jautājumus par novērošanas tehnoloģiju ētiku, finansējumu un demokrātisko uzraudzību ES, kā arī par ES pētniecības programmu atbildību pret personām, kas nav ES pilsoņi. Eiropas Parlamenta deputāts Patriks Breijers pieprasīja publiskot dokumentus, kas saistīti ar iBorderCtrl, ES finansētu pētniecības projektu, kura mērķis bija ieviest automatizētu melu atklāšanu uz ES robežām. Breijeram sākotnēji tika liegta iespēja iegūt informāciju par sistēmu sabiedriskās drošības un komerciālo interešu dēļ.

Savā spriedumā Vispārējā tiesa atbalstīja Breijera argumentu, secinot, ka sabiedrības interesēs ir demokrātiska uzraudzības nodrošināšana pār novērošanas un kontroles tehnoloģiju izstrādi. Turklāt šis svarīgais nolēmums apstiprina, ka sabiedrības interesēs ir apstrīdēt to, vai šādas tehnoloģijas ir vēlamas un vai to izstrāde vispār būtu jāfinansē no valsts līdzekļiem. iBorderCtrl ir tikai viena no daudzajām “inovācijām”, kas ieviesta, lai novērtētu risku un pieņemtu lēmumus robežkontroles un imigrācijas jomā.

Šādas sistēmas izmantošana var radīt daudzus iespējamus tiesību pārkāpumus, tai skaitā Hartas 18. pantā noteikto patvēruma tiesību pārkāpumu. Ja robežsargs, kuru informē MI melu detektora modulis, kas maldīgi norāda, ka personas mikroizteiksmes ir negodīgas, atsaka ceļotājam ieceļošanu, tas var pārkāpt šīs personas tiesības. Eiropas datu aizsardzības

⁵³ [Gallagher R., Jona L. \(26 July, 2019\). We tested Europe's new lie detector for travelers – and immediately triggered a false positive. The Intercept.](#)

⁵⁴ Vispārējās Tiesas 2021. gada 15. decembra spriedums lietā T-158/19 Patrick Breyer pret Eiropas Pētniecības izpildaģentūru, ECLI:EU:T:2021:902

bijušais uzraudzītājs Džovanni Butarelli (Giovanni *Buttarelli*) ir vērsis uzmanību, ka minētā sistēma var diskriminēt cilvēkus viņu etniskās izcelsmes dēļ. Šāda veida sistēmas būtiski apdraud cilvēktiesības, jo īpaši diskriminācijas aizlieguma principa ievērošanu.⁵⁵

Eiropas Digitālo tiesību asociācija (European Digital Rights, EDRI – angļu val.) norāda, ka MI testēšanai uz Eiropas robežām, lai it kā atklātu melus, izmantojot imigrācijas lietotnes, vai maldināšanu par angļu valodas testiem, izmantojot balss analīzi, trūkst ticama zinātniska pamatojuma. ES migrācijas politiku arvien vairāk atbalsta MI sistēmas, piemēram, sejas atpazīšana, algoritmiskās profilēšanas un prognozēšanas rīki, kas paredzēti migrācijas pārvaldības procesos, tostarp piespiedu izraidīšanai. Šie izmantošanas gadījumi var pārkāpt datu aizsardzības tiesības, tiesības uz privātumu, tiesības uz nediskrimināciju un vairākus starptautisko migrācijas tiesību principus, tostarp tiesības meklēt patvērumu.⁵⁶

MI sistēmu izmantošana, lai palīdzētu izskatīt patvēruma, vīzu vai uzturēšanās atļauju pieteikumus

Mākslīgā intelekta akts paredz, ka augstu risku rada MI sistēmas, kuras izmanto, lai palīdzētu kompetentām publiskajām iestādēm izskatīt patvēruma, vīzu vai uzturēšanās atļauju pieteikumus, kā arī saistītās sūdzības, nosakot, vai fiziskās personas, kuras iesniegušas pieteikumu uz kādu no minētajiem statusiem, atbilst kritērijiem, tostarp veicot saistītu pierādījumu patiesuma novērtēšanu (III pielikuma 7. punkta c) apakšpunkts).

Viens no piemēriem, kas parāda šādu sistēmu diskriminācijas risku ir **Nīderlandes profilēšanas sistēma** īstermiņa vīzu pieteikumu izvērtēšanai. Lai izvērtētu personu, kuras vēlas ieceļot Nīderlandē un Šengenas zonā, īstermiņa vīzu pieteikumus, tika izmantota profilēšanas sistēma, kuras pamatā bija tādi kritēriji kā tautība, dzimums un vecums. Ja sistēma klasificēja pieteikuma iesniedzēju kā augsta riska personu, iestādes veica papildu izmeklēšanu, kas bieži vien noveda pie pieteikuma izskatīšanas kavēšanas un diskriminējošas

⁵⁵ Sk. [Barkane, I. \(2023\). Cilvēktiesību nozīme mākslīgā intelekta laikmetā. Privātums, Datu aizsardzība un regulējums masveida novērošanas novēršanai. Rīga: LU Akadēmiskais apgāds.](#)

⁵⁶ [EDRI. \(12 January, 2021\). Re: Open letter: Civil society call for the introduction of red lines in the upcoming European Commission proposal on Artificial Intelligence.](#)

neobjektivitātes. Piemēram, vīzu iegūšana Nīderlandē dzīvojošajiem Marokas un Surinamas pilsoņu ģimenes locekļiem bija sarežģīta.⁵⁷

Labi zināma problēma, kas saistīta ar MI sistēmām un automatizētu riska novērtējumu, ir diskriminācijas saglabāšanas vai reproducēšanas iespēja. Ja šāda veida sistēmas tiek apmācītas ar vēsturiskiem datiem, piemēram, neautomatizētiem lēmumiem, ko vīzu piešķiršanas procesā pieņem darbinieki, lai atpazītu potenciālos nelegālos imigrantus, pastāv risks atkārtot diskrimināciju, kas ir šo cilvēku lēmumu pamatā, kuri bieži vien balstās uz **etnisko un rasu** profilēšanu. Visbeidzot, šīs sistēmas var pieļaut kļūdas, kas var novest pie diskriminācijas, palielinot nepamatotu ieceļošanas atteikumus vai migrantu nepareizu riska klasifikācijas gadījumus.⁵⁸

Vēl viens līdzīgs piemērs ārpus Eiropas ir **Chinook sistēma**, ko izmanto Imigrācijas, bēgļu un pilsonības pārvalde **Kanādā**. Tā ir uz Microsoft Excel balstīta programmatūra, kas vienkāršo pagaidu uzturēšanās atļauju pieteikumu apstrādi, bieži vien pat pieteicējiem nezinot, ka viņi tiek pakļauti algoritmam. Sistēma organizē vīzu pieteikumus izklājlapās un atvieglo masveida atteikumus, pamatojoties uz standartizētiem filtriem, radot nopietnas bažas par procesuālo taisnīgumu.

Lai gan valdība apgalvo, ka Chinook sistēmas netiek izmantota automatizēta lēmumu pieņemšanai, tās izmantošana praksē ir niansētāka. Amatpersonas izmanto automatizētas šķirošanas un apkopošanas funkcijas, neveicot daudzu pieteikumu individuālu izskatīšanu. Praksē tas ļauj pieņemt masveida atteikumus bez detalizēta pamatojuma, kas nesamērīgi ietekmē pretendentes no valstīm ar augstākiem atteikumu rādītājiem, galvenokārt no globālajiem dienvidiem.

Kanādas Federālā tiesa ir pārbaudījusi vairākus lēmumus, kas pieņemti, izmantojot Chinook sistēmu. Tā ir vērsusi uzmanību uz personalizētas analīzes trūkumu, kritizējusi vispārīgās

⁵⁷ [Pina M.C. \(2025\). AI in the context of border management, migration and asylum in the EU: technological innovation vs. fundamental rights of migrants in the AI Act.](#)

⁵⁸ Turpat.

atteikuma vēstules, kas ģenerētas, izmantojot Chinook, un uzsvērusi, ka lēmumiem jāatspoguļo individuāli novērtējumi un saprotams pamatojums.⁵⁹

Valodas analīze, lai noteiktu patvēruma meklētāju izcelsmes valsti

Eiropas Parlamenta 2025. gadā publicētajā ziņojumā “Mākslīgais intelekts patvēruma procedūrās” minēti MI balstītas valodas analīzes izmantošanas piemēri, lai noteiktu patvēruma meklētāju izcelsmes valsti.⁶⁰ Saskaņā ar Eiropas Savienības Patvēruma aģentūras (EUAA) 2022. gada ziņojumu septiņas ES dalībvalstis (Dānija, Vācija, Nīderlande, Austrija, Rumānija, Somija un Zviedrija) un Šveice ir izmantojušas valodas analīzes rīku patvēruma meklētāju izcelsmes valsts noteikšanai (angļu val. – Language assessment for the determination of origin (LADO)).⁶¹ Sešas citas dalībvalstis (Grieķija, Horvātija, Malta, Polija, Portugāle un Slovākija) apsver iespēju ieviest LADO tuvākajā nākotnē. LADO parasti tiek izmantots ar profesionālu valodnieku palīdzību. Ziņojumā tika atzīts, ka MI spēju attīstība, lai palīdzētu noteikt valodas un dialektu, rada jaunas iespējas LADO izmantošanā, tajā pašā laikā tajā ir arī norādīts uz dažu ieinteresēto personu bažām par potenciālām precizitātes un līdz ar to arī uzticamības problēmām, ko rada MI rīku izmantošana.

Kopš 2023. gada EUAA ir aktīvi iesaistījusies diskusijās ar dalībvalstīm par to, kā apvienot MI un cilvēka veiktu izvērtēšanu, lai uzlabotu rezultātu kvalitāti LADO jomā.⁶² Laikā no 2021. līdz 2022. gadam aģentūra piedalījās pilotprojektā, kurā tika testēts hibrīda valodas novērtēšanas modelis, kas apvienoja balss paraugu analīzi, ko veic mašīna, un cilvēku ekspertu analīzi.

Vācijā Federālais migrācijas un bēgļu birojs (BAMF) izmanto uz MI balstītu dialektu atpazīšanas sistēmu (Valodas biometrijas palīdzības sistēma — DIAS), lai atbalstītu lēmumu pieņemējus patvēruma meklētāju izcelsmes valsts un/vai reģiona identificēšanā.⁶³ Sistēma

⁵⁹ [Orakhelashvili, A. \(2024\). Real World Cases of Annex III AI Act applications that posed risks to fundamental rights.](#); [Emian, Y. \(2025.\) Digital Borders and Racial Codes in AI Migration Control.](#)

⁶⁰ [European Parliament. \(2025\). Artificial intelligence in asylum procedures in the EU.](#)

⁶¹ [Skat. European Union Agency for Asylum \(EUAA\). \(2022\). Study on Language Assessment for Determination of Origin of Applicants for International Protection.](#)

⁶² [Skat. European Union Agency for Asylum \(EUAA\). \(2023\). Strategy on Digital Innovation in Asylum Procedures and Reception Systems.](#)

⁶³ [Späth, E. \(2025\). AI Use in the Asylum Procedure in Germany: Exploring Perspectives with Refugees and Supporters on Assessment Criteria and Beyond. In: Ahrweiler, P. \(eds\) Participatory Artificial Intelligence in Public Social Services.](#)

tiek izmantota kopš 2017. gada noteiktiem arābu valodas dialektiem (ēģiptiešu, līča, irākiešu, levantu, magrebu), bet kopš 2022. gada — dari, persiešu un puštu valodai. DIAS analizē fonētiskos modeļus pieteikuma iesniedzēju runā, lai varbūtiski noteiktu viņu izcelsmes valsti vai reģionu. Novērtēšanai pretendentiem patvēruma procedūras sākumā tiek lūgts aprakstīt attēlu savā valodā, veicot telefona zvanu. Šaubu gadījumā DIAS var tikt izmantots atkārtoti intervijas laikā. Saskaņā ar BAMF datiem, DIAS ziņojums ir tikai papildinošs un neautomatizē identifikācijas procesu vai ticamības novērtējumu. Pretendentiem ir iespēja komentēt konstatējumus intervijas laikā.

DIAS ir daļa no integrētajiem identitātes pārvaldības rīkiem, kas ietver arī rīkus automātiskai sejas atpazīšanai, vārdu transliterācijai un datu ierīču analīzei. Vārdu transliterācijas rīks tiek izmantots, lai novērstu pareizrakstības kļūdas un standartizētu vārdu pareizrakstību, kas sākotnēji rakstīti ne latīņu alfabētā. Iestādes apgalvo, ka rīks arī palīdz identificēt pieteikuma iesniedzēju izcelsmes valsti un tādējādi novērtēt pieteikuma iesniedzēju apgalvojumu ticamību par viņu izcelsmes vietu. Saskaņā ar BAMF datiem, Vācija sadarbojas ar vairākām Eiropas valstīm, lai "plānotu valodas un dialektu atpazīšanas pilotprojektu, kurā tiks pārbaudīta runas ierakstu apmaiņa un analīze".

Itālijā iestādes ir testējušas MI rīku (S.I.N.D.A.C.A), lai automātiski transkribētu intervijas ar patvēruma meklētājiem. Saskaņā ar pakalpojumu sniedzēja teikto, "automātiskā transkripcija ņem vērā visus semantiskos aspektus, neizslēdzot dialektu, akcentu, svešvalodu terminoloģijas un spontānas runas atpazīšanu, ar precizitātes līmeni ne mazāk kā 95%".⁶⁴

Laikā no 2019. līdz 2021. gadam ES finansēja projektu Turcijas Republikā, kurā tika izmēģināta valodas analīzes sistēma, kas paredzēta pilsonības noteikšanas procedūru uzlabošanai, īpaši koncentrējoties uz uiguru un uzbeku pilsoņu atšķiršanu. Saskaņā ar AFAR ziņojumu Zviedrijas Migrācijas aģentūra testēja MI rīku valodas atpazīšanai, bet neturpināja tā ieviešanu.⁶⁵ Vairākas citas dalībvalstis plānoja testēt dialektu atpazīšanas tehnoloģijas. Ungārija veica sākotnējos pētījumus par dialektu atpazīšanas tehnoloģijas izmantošanu patvēruma iestādēs. Horvātija ir paziņojusi par plāniem nākotnē izmantot valodas

⁶⁴ [European Parliament. \(2025\). Artificial intelligence in asylum procedures in the EU.](#)

⁶⁵ [Ozkul, D. \(2023\). Automating Immigration and Asylum: The Uses of New Technologies in Migration and Asylum Governance in Europe.](#)

identifikāciju kā daļu no savas patvēruma procedūras, lai noteiktu pretendentu izcelsmes valsti.⁶⁶

Nīderlandē imigrācijas iestādes izmanto teksta izpētes rīku (Case Matcher), kas atsijā patvēruma pieteikumus, lai identificētu pieteikumus, kas iesniegti uz līdzīga pamata.⁶⁷

Lietotnes mērķis ir palīdzēt samazināt laiku, kas lietu izskatītājiem jāpavada, izskatot jaunas lietas, un nodrošināt konsekveni visos lēmumos. Šis rīks varētu sniegt lēmumu pieņēmējiem dziļāku izpratni par izplatītajiem riskiem pretendenta izcelsmes valstī. Tomēr tas varētu ietekmēt arī pieteikumu ticamības novērtējumu, jo pretendenti ar līdzīgiem stāstiem varētu tikt turēti aizdomās par melošanu.

Lai gan MI sistēmas pēdējās desmitgades laikā ir ievērojami attīstījušās, tās joprojām var būt neprecīzas un neobjektīvas. Neprecīzi un/vai neobjektīvi MI novērtējumi vai ieteikumi par svarīgiem patvēruma pieteikumu aspektiem var izraisīt nopietnus pamattiesību pārkāpumus. Neobjektīvi algoritmi var izraisīt diskrimināciju pret **sievietēm, etniskajām minoritātēm, vecāka gadagājuma cilvēkiem un citām grupām**.⁶⁸

Tā kā MI sistēmas tiek izstrādātas, izmantojot liela apjoma datus, šo datu kvalitāte ir būtiska, lai garantētu šo sistēmu precizitāti un uzticamību. Slikta kvalitātes apmācības dati noved pie sliktiem rezultātiem. ES Pamattiesību aģentūra ir uzsvērusi, ka kļūdas datu analīzē vai interpretācijā var izraisīt nepareizus secinājumus par pieteikuma iesniedzēja izcelsmi, kas noved pie negodīgiem lēmumiem ar potenciāli dzīvībai bīstamām sekām attiecīgajai personai.⁶⁹ Pietiekamu un reprezentatīvu datu trūkums var radīt papildu problēmas. Piemēram, dialektu atpazīšanas rīks, kas apmācīts ar datu kopu, kurā nav iekļauti konkrēta reģiona dialekti, var nepareizi atpazīt šo reģionu dialektus. Valodu gadījumā apmācības dati var kļūt arī nepilnīgi, jo valodas un dialekti laika gaitā attīstās.

Lai gan MI dažkārt tiek slavēts kā instruments neobjektivitātes un diskriminācijas mazināšanai (samazinot cilvēka rīcības brīvību), pastāv bažas, ka tas varētu izraisīt un palielināt diskrimināciju patvēruma procedūrās. MI algoritmi, kas apmācīti ar datiem no

⁶⁶ [European Parliament. \(2025\). Artificial intelligence in asylum procedures in the EU.](#)

⁶⁷ Turpat.

⁶⁸ [European Union Agency for Fundamental Rights \(FRA\). \(2019\). Data quality and artificial intelligence – mitigating bias and error to protect fundamental rights.](#)

⁶⁹ Turpat.

iepriekšējiem lēmumiem, varētu nākotnē radīt aizspriedumus, kas bija iestrādāti šajos lēmumos. Tas varētu radīt kaitīgas atgriezeniskās saites cilpas, kur algoritma radītie izejas dati, kas satur aizspriedumus un diskrimināciju, tiek atgriezti sistēmā kā ievades vai apmācības dati. Piemēram, pētījumi par MI lietojumprogrammām, ko izmanto migrācijas procedūrās Austrālijā un Kanādā, atklāja, ka šīs sistēmas pastiprina jau esošās atšķirības imigrācijas rezultātos.⁷⁰

Vēl viena no problēmām ir tā, ka MI sistēmas varētu paļauties uz starpniekservera datiem un maldīgiem pieņēmumiem par šo datu atbilstību. Piemēram, dialektu atpazīšanas rīki tiek izmantoti, lai noskaidrotu pretendentu izcelsmes valsti un tautību. Tomēr tie var neņemt vērā indivīda reģionālo socializāciju ne viņa izcelsmes valstī, ne ceļojuma laikā. Turklāt valodu prasmes nav labs tautības rādītājs.

Problēmas var rasties arī tad, ja dati nav vienmērīgi sadalīti, piemēram, ja patvēruma pieteikumi pārsvarā attiecas uz noteiktām tautībām un konkrētiem vajāšanas veidiem. Tas var radīt grupas nelīdzsvarotības problēmu, kur modelī ir pārāk pārstāvētas vairākuma grupas īpašības (piemēram, Sīrijas patvēruma meklētāji, kas bēg no pilsoņu kara), kas noved pie neobjektīvām prognozēm un samazinātas veikspējas attiecībā uz minoritāšu grupām (piemēram, pretendenti no citiem reģioniem, kas meklē aizsardzību citu iemeslu dēļ).

Problemātiski var būt arī tas, kā lietotāji interpretē un izmanto MI sistēmu sniegtos rezultātus. MI sistēmas parasti sniedz varbūtības apgalvojumus, kas nozīmē, ka to rezultāti nevar nozīmēt noteiktību. Tām ir arī kļūdu līmenis, kas jāņem vērā. Piemēram, ziņots, ka Vācijā izmantotā DIAS kļūdu līmenis arābu dialektiem ir no 15 % līdz 20 %, bet persiešu dialektiem — 27 %. Vārdu transliterācijas rīka veiksmes līmenis pretendentiem no Magribas valstīm ir tikai 35 %. Starp pārbaudītajiem arābu vārdiem 39 % atsauču uz izcelsmes valsti nebija pārbaudāmas. Pastāv risks, ka neprecīzi vai nepārliciecināmi rezultāti var pastiprināt aizdomas par pretendentu apgalvojumiem. Šo algoritmu necaurredzamība apgrūtina problēmas identificēšanu un rezultātu apstrīdēšanu, kas rada labvēlīgu vidi algoritmiskai diskriminācijai.

⁷⁰ [European Parliament. \(2025\). Artificial intelligence in asylum procedures in the EU.](#)

Ņemot vērā MI sistēmu, kas novērtē ieceļotājus, iespējamo augsto aizskārumu un zinātniskā pamatojuma trūkumu, pētnieki ir paiduši uzskatu, ka automatizētas riska novērtēšanas aizliegumam nevajadzētu attiekties tikai uz gadījumiem, kas saistīti ar indivīda noziedzīga nodarījuma izdarīšanas riska prognozēšanu. Faktiski, lai gan automatizēta drošības, nelegālās migrācijas vai veselības risku novērtēšana migrācijas kontekstā tiek uzskatīta par augsta riska un tai ir jāatbilst noteiktām prasībām, bažas rada tas, ka šāda prakse ir atļauta kontekstā, ko raksturo dziļi iesakņojusies etniskā, rasu un dzimumu diskriminācija un migrantu paaugstināta neaizsargātība.⁷¹

Līdzās MI sistēmām, kas palīdz izvērtēt risku, ko rada ieceļotāji, tiek ieteikts augsta riska sistēmu sarakstā iekļaut MI sistēmas, kuru mērķis ir prognozēt migrācijas tendences un robežšķērsošanu, lai gan tās šķiet mazāk iejaucas pamattiesībās. Piemēram, Eiropas Patvēruma atbalsta birojs (European Asylum Support Office (EASO) – kopš 2022. gada – izstrādāja Agrīnās brīdināšanas un sagatavotības sistēmu, kas paredzēta migrācijas plūsmu prognozēšanai ES teritorijās. Šī sistēma balstās uz tādiem datu avotiem kā GDELT (informācija par notikumiem pēc izcelsmes valsts), Google Trends (iknedēļas tiešsaistes meklēšanas tendences pēc izcelsmes valsts), Frontex (ikmēneša nelegālās robežšķērsošanas atklāšana) un iekšējiem datiem par patvēruma pieteikumu skaitu un atzīšanas rādītājiem ES dalībvalstīs. Algoritms cenšas paredzēt, kuri notikumi izraisīs liela mēroga pārvietošanos, un aplēst turpmāko patvēruma pieteikumu skaitu ES.⁷²

No vienas puses, prognozējot migrantu ierašanos, šīs sistēmas var nodrošināt efektīvu sagatavošanos cilvēku ierašanās brīdim un ļaut pārdalīt resursus atbilstoši uzņemšanas vajadzībām. No otras puses, tās var veicināt preventīvus pasākumus, lai kavētu migrācijas kustību, veicot pasākumus, kas kavē migrantu un patvēruma meklētāju piekļuvi valsts teritorijai.⁷³ MI var kļūt par vēl vienu politisku instrumentu, ko izmanto, lai pastiprinātu valstu prakses, kuru mērķis ir ierobežot starptautisko migrāciju un neļaut patvēruma

⁷¹ [Pina, M.C. \(2025\). AI in the context of border management, migration and asylum in the EU: technological innovation vs. fundamental rights of migrants in the AI Act.](#)

⁷² [Ozkul, D. \(2023\). Automating Immigration and Asylum: The Uses of New Technologies in Migration and Asylum Governance in Europe.](#)

⁷³ [Brouwer, E. \(2024\). EU's AI Act and migration control. Shortcomings in safeguarding fundamental rights, VerfBlog.](#)

meklētājiem sasniegt to teritorijas.⁷⁴ Līdz ar to šīm sistēmām ir jāpiemēro stingrs regulējums.

⁷⁴ [Beduschi A. \(2020\). International migration management in the age of artificial intelligence, *Migration Studies*, 9\(3\), 576-596.](#)

6. Tiesvedība un demokrātijas procesi

Tiesvedības un demokrātijas procesu jomā Mākslīgā intelekta akta III pielikuma 8. punkts nosaka divu veidu sistēmas, kuru izmantošana rada augstu risku. Pirmkārt, augstu risku rada MI sistēmas, ko paredzēts izmantot tiesu iestādē vai tās vārdā, lai palīdzētu tiesu iestādei izpētīt un interpretēt faktus vai tiesību normas un piemērot tiesību normas konkrētam faktu kopumam, vai izmantot līdzīgā veidā strīdu alternatīvā izšķiršanā. Par augsta riska sistēmām nav uzskatāmas MI sistēmas, kuras paredzētas tikai tādu administratīvu papilddarbību veikšanai, kam nav ietekmes uz faktisko tiesvedību atsevišķās lietās, piemēram, tiesas nolēmumu, dokumentu vai datu anonimizācijai vai pseidonimizācijai, saziņai starp darbiniekiem un administratīvu uzdevumu veikšanai (61. apsvērums).

Otrkārt, par augsta riska sistēmām ir uzskatāmas MI sistēmas, ko paredzēts izmantot, lai ietekmētu vēlēšanu vai referenduma rezultātus vai fizisku personu uzvedību balsošanā, kad tās balso vēlēšanās vai referendumos. MI sistēmas, kas paredzētas izmantošanai tiesvedībā un demokrātiskajos procesos, ir klasificētas kā augsta riska sistēmas, ņemot vērā to potenciāli ievērojamo ietekmi uz demokrātiju, tiesiskumu, individuālajām brīvībām, kā arī tiesībām uz efektīvu tiesību aizsardzību un taisnīgu tiesu (Mākslīgā intelekta akta 61. apsvērums).

6.1. Demokrātijas procesi un MI

Attiecībā uz MI sistēmām, kas var ietekmēt demokrātiskos procesus, Mākslīgā intelekta akts skaidro, ka neskarot Eiropas Parlamenta un Padomes Regulā (ES) 2024/900 par politiskās reklāmas pārredzamību un mērķorientēšanu⁷⁵ paredzētos noteikumus un lai novērstu riskus, kas saistīti ar nepamatotu ārēju iejaukšanos Hartas 39. pantā nostiprinātajās tiesībās balsot un nelabvēlīgu ietekmi uz demokrātiju un tiesiskumu, MI sistēmas, ko paredzēts lietot, lai ietekmētu vēlēšanu vai referenduma rezultātus vai fizisku personu balsošanas uzvedību, kad tās piedalās vēlēšanās vai referendumos, ir klasificējamas kā augsta riska MI sistēmas (62. apsvērums). Tajā pašā laikā par augsta riska MI sistēmām netiek uzskatītas sistēmas, kuru rezultātam fiziskas personas nav tieši pakļautas, piemēram, rīkiem, ko

⁷⁵ [Komisijas deleģētā regula \(ES\) 2023/2450 \(2023. gada 25. jūlijs\), ar ko papildina Eiropas Parlamenta un Padomes Direktīvu \(ES\) 2022/2557, izveidojot pamatpakalpojumu sarakstu. OV L, 2023/2450, 30.10.2023.](#)

izmanto, lai organizētu, optimizētu vai strukturētu politiskas kampaņas no administratīvā vai loģistikas viedokļa (Mākslīgā intelekta akta III pielikuma 8. punkta b) apakšpunkts, 62. apsvērums).

Ir jānošķir augsta riska MI sistēmas, ko izmanto tiesvedības un demokrātijas procesu jomā, no aizliegtas MI prakses, kas var apdraudēt demokrātijas procesus. Proti, saskaņā ar Mākslīgā intelekta aktu ir aizliegta tādas MI sistēmas laišana tirgū, nodošana ekspluatācijā vai lietošana, kas izmanto cilvēka neapzinātus subliminālus paņēmienus vai tīši manipulatīvus vai maldinošus paņēmienus, kuru mērķis ir būtiski iespaidot personas vai personu grupas uzvedību vai kuriem ir tādas sekas, kas ievērojami pasliktina viņu spēju pieņemt informētu lēmumu, tādējādi liekot tām pieņemt lēmumu, ko tās citādi nebūtu pieņēmušas, tādā veidā, kas šai personai, citai personai vai personu grupai rada vai pamatoti ticami radīs būtisku kaitējumu (Mākslīgā intelekta akta 5. panta 1. punkta a) apakšpunkts). Mākslīgā intelekta akts aizliedz arī tādas MI sistēmas lietošanu, kura izmanto kādu no kādas fiziskas personas vai konkrētas personu grupas neaizsargātības iezīmēm viņu vecuma, invaliditātes vai konkrētas sociālās vai ekonomiskās situācijas dēļ un kuras mērķis vai sekas ir minētās personas vai minētajai grupai piederīgas personas uzvedības būtiska iespaidošana tādā veidā, kas attiecīgajai personai vai citai personai rada vai ir saprātīgi iespējams, ka radīs būtisku kaitējumu (Mākslīgā intelekta akta 5. panta 1. punkta b) apakšpunkts).

Mākslīgā intelekta aktā “maldinoši paņēmieni”, kas minēti 5. panta 1. punkta a) apakšpunktā, nav definēti. Tomēr Mākslīgā intelekta akta 29. apsvērumā ir skaidrots, ka “maldinoši paņēmieni” ir paņēmieni, kas grauj vai vājina personas patstāvību, lēmumu pieņemšanu vai brīvu izvēli veidā, kuru persona neapzinās vai, ja tomēr apzinās, tā joprojām var tikt maldināta vai kuru tā nespēj kontrolēt vai kuram pretoties. MI sistēmu “maldinošie paņēmieni” būtu jāsaprot kā tādi, kas ietver nepatiesas vai maldinošas informācijas sniegšanu, lai maldinātu personas vai iespaidotu viņu uzvedību tādā veidā, kas apdraud viņu patstāvību, lēmumu pieņemšanu un brīvu izvēli vai rada šādas sekas.

Mākslīgā intelekta akts ir vērsts uz pieaugošām bažām par dziļviltojumu. Dziļviltojums ir MI ģenerēts vai manipulēts attēls, audio vai video saturs, kurš atgādina reālas personas, priekšmetus, vietas, vienības vai notikumus un maldina personu, ka attiecīgais saturs ir autentisks vai patiess (Mākslīgā intelekta akta 3. panta 60. apakšpunkts). Mākslīgā intelekta

akts regulē dziļviltojumus, pamatojoties uz to radītā riska līmeni. Lai gan Mākslīgā intelekta akts MI sistēmas, kuras lieto, lai radītu dziļviltojumu, klasificē kā ierobežota riska sistēmas, paredzot tām pārredzamības pienākumu (50. pants), MI sistēmas, ko izmanto, lai ietekmētu vēlēšanas, referendumu vai fizisku personu balsošanas uzvedību, tiek uzskatītas par augsta riska sistēmām saskaņā ar III pielikuma 8. punkta b) apakšpunktu, kā arī to izmantošana var būt aizliegta MI prakse saskaņā ar Mākslīgā intelekta akta 5. pantu.

EK Pamatnostādnes norāda, ka pastāv mijiedarbību starp šo Mākslīgā intelekta akta 5. panta 1. punkta a) apakšpunktā noteikto aizliegumu un Mākslīgā intelekta aktā noteikto pārredzamības pienākumu – 50. panta 4. punktā uzturētāja pienākumu marķēt “dziļviltojumus” un dažas MI ģenerētas teksta publikācijas par jautājumiem, kas ir sabiedrības interesēs, kā arī nodrošinātāja pienākumu nodrošināt, ka MI sistēmas, kas mijiedarbojas ar cilvēkiem, ir izstrādātas tā, lai informētu cilvēkus par to, ka viņi mijiedarbojas ar MI, nevis cilvēku. Šāda redzama informācijas atklāšana ir riska mazināšanas pasākums, kas būtu jānodrošina arī projektējot MI sistēmas, tai skaitā ar tehniskiem pasākumiem, kas ļauj atklāt MI radītu un manipulētu saturu. Pamanāms “dziļviltojumu” un sarunbotu marķējums samazina maldināšanas risku, kas varētu rasties, kad MI radītais saturs tiek izplatīts sabiedrībā, kā arī samazina risku, ka tiks radīta kaitīga ietekme uz personas viedokļa un pārliecības veidošanos un uzvedību. Turpretī Mākslīgā intelekta akta 5. panta 1. punkta a) apakšpunktā noteiktajam aizliegumam ir daudz ierobežotāka darbības joma. Tas var attiekties, piemēram, uz gadījumiem, kad sarunbots vai maldinošs MI radīts saturs sniedz nepatiesu vai maldinošu informāciju veidos, kuru mērķis vai rezultāts ir personu maldināšana un uzvedības iespaidošana, kas nebūtu notikusi, ja viņi nebūtu bijuši pakļauti mijiedarbībai ar MI sistēmu vai maldinošu MI radīto saturu, jo īpaši, ja tas netiek redzami atklāts.⁷⁶

Mākslīgā intelekta akts nosaka pārredzamības pienākumu, liekot dziļviltojumu veidotājiem informēt sabiedrību par sava darba mākslīgo raksturu. Papildus tehniskajiem risinājumiem, ko pielieto MI sistēmas nodrošinātāji, uzturētāji, kuri MI sistēmu lieto, lai ģenerētu vai manipulētu tādu attēlu, audio vai video saturu, kas ievērojami atgādina reālas personas,

⁷⁶ [Eiropas Komisija. \(2025\). Komisijas paziņojums. Komisijas Pamatnostādnes par aizliegtu mākslīgā intelekta praksi, kas noteikta Regulā \(ES\) 2024/1690 \(MI aktā\).](#)

objektus, vietas, subjektus vai notikumus un kas personai maldinoši šķistu autentisks vai īsts (dziļviltrojumi), būtu arī skaidri un atšķiramā veidā jāizpauž, ka saturs ir mākslīgi radīts vai manipulēts, attiecīgi marķējot MI radīto iznākumu un izpaužot tā mākslīgo izcelsmi (134. apsvērums). Saskaņā ar Mākslīgā intelekta akta 50. panta 4. daļu MI sistēmas uzturētāji, kura ģenerē vai manipulē attēla, audio vai video saturu, kas ir dziļviltrojums, izpauž, ka saturs ir mākslīgi ģenerēts vai manipulēts. Līdzīgi arī tādas MI sistēmas uzturētāji, kura ģenerē vai manipulē tekstu, ko publicē nolūkā informēt sabiedrību par sabiedrības interešu jautājumiem, izpauž, ka teksts ir mākslīgi ģenerēts vai manipulēts. Šo pienākumu nepiemēro, ja lietošana ir likumiski atļauta noziedzīga nodarījuma atklāšanai, novēršanai, izmeklēšanai vai kriminālvajāšanai vai ja MI ģenerētu saturu ir pārskatījis cilvēks vai tam ir veikta redakcionāla kontrole un ja fiziska vai juridiska persona ir redakcionāli atbildīga par satura publikāciju.

Lai gan principā Mākslīgā intelekta akta 50. pantā noteikto pārredzamības nodrošināšanas pienākumu mērķis ir līdz minimumam samazināt dziļviltrojumu un sarunbotu manipulatīvo ietekmi, tomēr var būt gadījumi un situācijas, kad neatkarīgi no informatīvajiem paziņojumiem šie maldinošie paņēmieni joprojām var būtiski ietekmēt personas un iespaidot viņu uzvedību tādā mērā, kas apdraud personu individuālo patstāvību un informētu lēmumu pieņemšanu, tāpēc tos nedrīkst ļaunprātīgi izmantot dezinformācijas un manipulācijas nolūkos. Dažos gadījumos uz tiem joprojām var attiekties Mākslīgā intelekta akta 5. panta 1. punkta a) apakšpunktā noteiktais aizliegums, ja ir izpildīti visi pārējie aizlieguma nosacījumi (tai skaitā attiecībā uz būtisku kaitējumu).⁷⁷

EK Pamatnostādnes skaidro, ka Mākslīgā intelekta akta 5.panta 1.punkta a) un b) apakšpunktā noteiktie aizliegumi neskar citus ES tiesību aktus un papildina tos. Tāda pati prakse, uz kuru attiecas Mākslīgā intelekta akta 5. panta 1. punkta a) vai b) apakšpunktā noteiktais aizliegums, var būt arī citu ES tiesību aktu pārkāpums, un uz to var attiecināt izpildi gan saskaņā ar Mākslīgā intelekta aktu, gan citiem tiesību aktiem. Tas ir svarīgi, jo šo tiesību aktu dažādo noteikumu mērķis ir aizsargāt dažādas intereses, un tiem ir atšķirīgi mērķi, darbības joma un adresāti. Tas nodrošina visaptverošu regulatīvo pieeju, kas aizsargā

⁷⁷ [Eiropas Komisija. \(2025\). Komisijas paziņojums. Komisijas Pamatnostādnes par aizliegtu mākslīgā intelekta praksi, kas noteikta Regulā \(ES\) 2024/1690 \(MI aktā\).](#)

personas un personu grupas no kaitīgas MI izmantošanu un manipulācijām un nodrošina drošus un uzticamus, MI iespējos pakalpojumus un produktus ES.⁷⁸

Mākslīgā intelekta aktā noteiktie aizliegumi atbilst arī ES datu aizsardzības tiesību aktiem, kuru mērķis ir aizsargāt datu subjektu personas datus, viņu pamattiesības un patstāvību. Lielāka personas datu apjoma pieejamība un lielākas iespējas apstrādāt šos datus ar MI sistēmām palielina kaitīgas manipulatīvas, maldinošas vai ekspluatatīvas prakses risku, piemēram, tādu, kas ir Mākslīgā intelekta akta 5. panta 1. punkta a) un b) apakšpunkta darbības jomā. Šajā kontekstā atbilstība datu aizsardzības noteikumiem par pārredzamību, datu minimizēšanu, taisnīgumu un likumību, piemēram, attiecībā uz personalizētu profilēšanu un reklāmu, kuras pamatā ir pakalpojuma lietotāju dati, var palīdzēt izvairīties no kaitējošas personalizētas manipulācijas un izmantošanas. Līdzīgi Mākslīgā intelekta akts neietekmē praksi, kas ir aizliegta ar ES diskriminācijas novēršanas tiesību aktiem.⁷⁹

Mākslīgā intelekta akta 5. panta 1. punkta a) un b) apakšpunktā noteiktie aizliegumi arī papildina Regulu (ES) 2022/2065 (Digitālo pakalpojumu aktu)⁸⁰, kas reglamentē tiešsaistes starpniecības pakalpojumus, piemēram, tiešsaistes platformas un meklētājprogrammas, un nodrošina pārredzamību un pārskatatbildību minēto pakalpojumu sniegšanā. Proti, saskaņā ar Digitālo pakalpojumu akta 25. panta 1. punktu lietotāja saskarnē nedrīkst izmantot tumšos modeļus, lai nodrošinātu, ka tiešsaistes platformu nodrošinātāji nemaldina vai nepiespiež lietotājus veikt darbības, kas var neatbilst to patiesajiem nodomiem. Šādi tumši modeļi būtu jāuzskata par manipulatīvu vai maldinošu paņēmienu piemēru Mākslīgā intelekta akta 5. panta 1. punkta a) apakšpunkta nozīmē, ja tie, visticamāk, radīs būtisku kaitējumu. Digitālo pakalpojumu akts arī nosaka tiešsaistes platformu nodrošinātāju pienākumus nodrošināt pārredzamību reklāmas jomā (26. un 38. pants attiecībā uz ļoti lielām tiešsaistes platformām vai ļoti lielām meklētājprogrammām), ieteikumu sistēmu izmantošanu (27. pants) un nepilngadīgo aizsardzību (28. pants). Turklāt, ja tiešsaistes platforma vai meklētājprogramma ir klasificēta kā ļoti liela tiešsaistes platforma vai ļoti liela meklētājprogramma, minētā norādītā pakalpojuma nodrošinātājam ir papildu pienākums

⁷⁸ Turpat.

⁷⁹ Turpat.

⁸⁰ [Eiropas Parlamenta un Padomes Regula \(ES\) 2022/2065 \(2022. gada 19. oktobris\) par digitālo pakalpojumu vienoto tirgu un ar ko groza Direktīvu 2000/31/EK \(Digitālo pakalpojumu akts, OV L, 227, 27.10.2022.](#)

novērtēt un mazināt sistēmisko risku, kas izriet no tā pakalpojuma un ar to saistīto sistēmu, tai skaitā algoritmiskās sistēmas, izstrādes vai darbības (34. un 35. pants). Veicot riska novērtējumus, ļoti lielu tiešsaistes platformu un ļoti lielu tiešsaistes meklētājprogrammu nodrošinātājiem būtu jāapsver, kā to ieteikumu sistēmas, reklāma, satura moderācija un visas citas attiecīgās algoritmiskās sistēmas ietekmē šādus sistēmiskos riskus. Šādos riska novērtējumos būtu arī jāanalizē, kā sistēmisko risku cita starpā ietekmē tīša manipulēšana ar pakalpojumu un tā automatizēta izmantošana (34. panta 2. punkts un 83. apsvērums). Tomēr Mākslīgā intelekta akta 5.panta 1.punkta a) un b) apakšpunkta darbības joma aptver plašu citu scenāriju klāstu (piemēram, sarunbotus, MI iespējos pakalpojumus un produktus), kurus var piedāvāt vai izmantot citi dalībnieki, kas nav starpniecības pakalpojumu sniedzēji.

Mākslīgā intelekta akts arī papildina Regulas (ES) 2024/900 par politiskās reklāmas pārredzamību un mērķorientēšanu⁸¹ noteikumus, tai skaitā pārredzamības un saistītos pienācīgas rūpības pienākumus, lai rādītu politiskās reklāmas un sniegtu saistītos pakalpojumus, kā arī noteikumus par mērķorientēšanas un reklāmas rādīšanas paņēmieni izmantošanu saistībā ar tiešsaistes politisko reklāmu. Šī regula aizliedz profilēšanu, ko veic, pamatojoties uz konkrētām personas datu kategorijām saistībā ar tiešsaistes politisko reklāmu un mērķtiecīgu tādu personu uzrunāšanu, kas ir vismaz vienu gadu jaunākas par valstī noteikto balsošanas vecumu. Turklāt mērķtiecīga rīcība un reklāmas rādīšanas metodes tiešsaistes politiskās reklāmas kontekstā var izmantot tikai tad, ja to pamatā ir personas dati, kas savākti no datu subjektiem un ar viņu nepārprotamu piekrišanu. Iepriekš minētā regula paredz arī papildu pārredzamības prasības, t. i., politiskās reklāmas atklāšanu, šādu paņēmieni un galveno parametru izmantošanas aprakstu un papildu informāciju par attiecīgo loģiku, arī par MI sistēmu izmantošanu. Mērķtiecīga politiskā reklāma, kuras pamatā ir personas datu apstrāde, saskaņā ar minēto regulu palīdzēs nodrošināt, ka vēlētāju profilēšana, politisko reklāmu mērķorientēšana un reklāmas rādīšana notiek likumīgas pārliecināšanas robežās.

⁸¹ [Eiropas Parlamenta un Padomes Regula \(ES\) 2024/900 \(2024. gada 13. marts\) par politiskās reklāmas pārredzamību un mērķorientēšanu, OV L, 2024/900, 20.3.2024.](#)

MI, tostarp ģeneratīvais MI, var atstāt būtisku ietekmi uz sabiedrību un demokrātiju kopumā, jo tas palielina informācijas manipulācijas, dezinformācijas un viltus ziņu risku. MI ģeneratīvie rīki dažu sekunžu laikā var radīt sintētiskos attēlus, audio un video, kā arī mērķtiecīgi atlasīt konkrētas auditorijas grupas. MI tehnoloģijas var izmantot politiķi, politiskie aktīvisti, kas vēlas ietekmēt vēlēšanu rezultātus.

Dažas politiskās kampaņas var saskarties ar MI sistēmām raksturīgiem aizspriedumiem. Ņemot vērā, ka MI tiek apmācīts, izmantojot vēsturiskus, enciklopēdiskus vai publiskus datus, MI sistēmas absorbē visus aizspriedumus, kas pastāv šajos datos. Piemēram, MI, kas apmācīts, izmantojot datus no interneta, kuros cita starpā ir daudz rasistiska satura, visticamāk, ģenerēs rasistiskus vēstījumus. Politiskām kampaņām ir jābūt piesardzīgām attiecībā uz MI sistēmās iestrādātajām neobjektivitātēm, kas ietekmē arī to vēstījumus.⁸²

Lai cīnītos ar šiem draudiem ne tikai ES, bet arī nacionālajā līmenī, tiek pilnveidots tiesiskais regulējums.

Lai regulētu MI izmantošanu priekšvēlēšanu aģitācijā, 2024. gadā Saeima pieņēma grozījumus Priekšvēlēšanu aģitācijas likumā⁸³, kas paredz pienākumu informēt vēlētājus par MI izmantošanu aģitācijas materiālu sagatavošanā. Grozījumi papildina likumu ar noteikumiem par MI sistēmu izmantošanu priekšvēlēšanu aģitācijā, paredzot, ja priekšvēlēšanu aģitācijas periodā apmaksātā priekšvēlēšanu aģitācijā vai aģitācijas materiālā tiek izmantots MI sistēmas radīts personas atveidojums vai realitātei neatbilstošs notikums (attēls, audio vai video saturs), tas skaidri un nepārprotami norādāms (3.¹ pants). Grozījumi arī paredz aizliegumu lietot automatizētas sistēmas priekšvēlēšanu aģitācijas veikšanai, izmantojot sociālo mediju (tehnoloģisku platformu, kas ļauj veidot, publicēt un izplatīt informāciju publiskā saziņā, kā arī veidot kopienas un savstarpēji mijiedarboties uz šā satura pamata) viltotu vai anonīmu kontu profilus (18. panta sestā daļa).

Lai cīnītos ar dziļviltojumiem vēlēšanu procesā, arī Krimināllikumā tika izdarīti divi grozījumi. 2024. gada 22. maijā stājas spēkā Krimināllikuma grozījumi⁸⁴, ar kuriem tika ieviesta kriminālatbildība par vēlēšanu procesa ietekmēšanu, izmantojot dziļviltojuma

⁸² [LaChapelle C., Tucker C. \(2023\). Generative AI in Political Advertising. The Brennan Center for Justice.](#)

⁸³ Grozījumi Priekšvēlēšanu aģitācijas likumā. Pieņemts 24.10.2024. Latvijas Vēstnesis, 06.11.2024, Nr. 217.

⁸⁴ Grozījumi Krimināllikumā. Pieņemts 09.05.2024. Latvijas Vēstnesis, 21.05.2024., Nr. 97.

tehnoloģiju. Krimināllikums tika papildināts ar jaunu 90.¹ pantu “Vēlēšanu procesa ietekmēšana, izmantojot dziļviltojuma tehnoloģiju”, kas noteic kriminālatbildību par apzināti nepatiesas un diskreditējošas informācijas izgatavošanu vai izplatīšanu par politisko partiju, to apvienību vai Saeimas, pašvaldības domes vai Eiropas Parlamenta deputāta kandidātu, izmantojot dziļviltojuma tehnoloģiju, ja tas izdarīts priekšvēlēšanu aģitācijas periodā vai vēlēšanu dienā. 2024. gada 22. oktobrī stājās spēkā Krimināllikuma grozījumi⁸⁵, kas ievieš kriminālatbildību par apzināti nepatiesas informācijas izgatavošanu vai izplatīšanu, izmantojot dziļviltojuma tehnoloģiju, lai diskreditētu valsts amatpersonas amata kandidātu, kuru amatā ievēlē, ieceļ vai apstiprina Saeima. Proti, Krimināllikums tika papildināts ar jaunu 90.² pantu “Valsts amatpersonas ievēlēšanas, iecelšanas vai apstiprināšanas amatā procesa Saeimā ietekmēšana, izmantojot dziļviltojuma tehnoloģiju”, kas paredz kriminālatbildību par apzināti nepatiesas, diskreditējošas informācijas izgatavošanu vai izplatīšanu, izmantojot dziļviltojuma tehnoloģiju, lai apmelotu valsts amatpersonas amata kandidātu, kuru amatā ievēlē, ieceļ vai apstiprina Saeima, ja tas izdarīts likumā noteiktajā valsts amatpersonas ievēlēšanas, iecelšanas vai apstiprināšanas amatā procesa laikā.

6.2. Tiesvedība un MI

Līdzās MI sistēmām, kas var ietekmēt demokrātiskos procesus, augstu risku rada arī MI izmantošana tiesvedībā. Mākslīgā intelekta akts skaidro, ka MI sistēmas, ko paredzēts izmantot tiesu darbā, tiek klasificētas kā augsta riska sistēmas, jo īpaši, lai mazinātu potenciālu neobjektivitātes, kļūdu un nepārredzamības risku. MI rīku lietošana var atbalstīt tiesnešu lēmumu pieņemšanas pilnvaras vai tiesu neatkarību, bet tai nevajadzētu to aizstāt. Proti, galīgajai lēmumu pieņemšanai arī turpmāk jābūt cilvēka virzītai darbībai. (61. apsvērums).

Līdzās ES MI regulējumam, 2024. gada 17. maijā Eiropas Padome pieņēma Pamatkonvenciju par mākslīgo intelektu, cilvēktiesībām, demokrātiju un tiesiskumu (EP

⁸⁵ Grozījumi Krimināllikumā. Pieņemts 26.09.2024. Latvijas Vēstnesis, 08.10.2024., Nr. 196.

Konvencija)⁸⁶, ko 2024. gadā 5. septembrī Viļņā Tieslietu ministru konferences laikā parakstīja ES un vairākas valstis. EP Konvencija paredz valstīm vispārīgu pienākumu aizsargāt pamattiesības (4. pants). Tāpat kā vispārīgs pienākums valstīm ir noteikts demokrātisko procesu integritātes un tiesiskuma ievērošana (5. pants). Tas paredz, ka valstis pieņem vai patur spēkā pasākumus, kuru mērķis ir nodrošināt, ka MI sistēmas netiek izmantotas, lai grautu demokrātisko institūciju un procesu integritāti, neatkarību un efektivitāti, tostarp varas dalīšanas principu, tiesu iestāžu neatkarības ievērošanu un piekļuvi tiesu iestādēm, kā arī pasākumus, kuru mērķis ir aizsargāt tās demokrātiskos procesus saistībā ar darbībām MI sistēmu dzīves ciklā, tostarp indivīdu taisnīgu piekļuvi un dalību publiskās debatēs, kā arī viņu spēju brīvi veidot viedokli. EP Konvencija nosaka vairākus pamatprincipus, kas saistīti ar MI darbību: cilvēka cieņa un individuālā autonomija (7. pants), pārredzamība un uzraudzība (8. pants), atbildība (9. pants), vienlīdzība un nediskriminācija (10. pants), privātus un personas datu aizsardzību (11. pants), uzticamība (12. pants), drošas inovācijas (13. pants). Tā paredz arī valstīm pienākumu nodrošinātu pieejamu un efektīvu tiesiskās aizsardzības līdzekļu pieejamību cilvēktiesību pārkāpumu gadījumā, kas izriet no darbībām MI sistēmu dzīves ciklā (14. pants). Gadījumos, kad MI sistēma būtiski ietekmē cilvēktiesību īstenošanu, personām, kuras tā skar, ir jābūt pieejamām efektīvām procesuālām garantijām, aizsardzības pasākumiem un tiesībām saskaņā ar piemērojamajiem starptautiskajiem un nacionālajiem tiesību aktiem (15. pants). EP Konvencija arī nosaka pienākumu izvērtēt un novērst riskus un negatīvu ietekmi, ko rada MI sistēmas (16. pants).

Lai gan abiem instrumentiem ir kopīgi mērķi aizsargāt pamattiesības un veicināt atbildīgu MI izstrādi, tie atšķiras pēc to juridiskā rakstura, darbības jomas un izpildes mehānismiem. Lai gan Mākslīgā intelekta akts paredz svarīgus soļus augsta riska MI sistēmu regulēšanā, tā pieeja tiesu MI sistēmām ir jāpilnveido, lai nodrošinātu saderību ar tiesu neatkarības un konstitucionālās pārvaldības pamatprincipiem. Šajā kontekstā var minēt vairākus piemērus: MI rīki, kas analizē iepriekšējos lēmumus, lai prognozētu lietu iznākumu, var netieši radīt spiedienu uz tiesnešiem, lai tie pielāgotos statistikas modeļiem, kas varētu apdraudēt tiesu autonomijas konstitucionālo aizsardzību;

⁸⁶ [Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law. 05.09.2024.](#)

nevienlīdzīga piekļuve MI balstītai juridiskajai analīzei juristu vidū varētu izjaukt līdzsvaru starp sacīkstes principiem; tiesu neatkarības noteikumi aizliedz automatizētu iejaukšanos juridiskajā argumentācijā. Tādējādi, lai gan Mākslīgā intelekta akta mērķis ir risināt ar MI saistītos riskus tiesu procesos, tas nepietiekami ņem vērā izveidoto tiesu pārvaldības struktūru lomu un rada bažas par atbilstošiem tiesās izmantoto MI sistēmu uzraudzības mehānismiem. Tā kā MI tehnoloģijas turpina attīstīties un iekļūst dažādos sabiedrības aspektos, šie regulējošie ietvari, visticamāk, veidos MI izstrādes, ieviešanas un pārvaldības nākotni tieslietu jomā. Turpmākie izaicinājumi ir ievērojami, tostarp nepieciešamība pastāvīgi pielāgot noteikumus, lai tie neatpaliktu no tehnoloģiju attīstības, un saglabāt līdzsvaru starp inovācijām un tiesu lietotāju aizsardzību.⁸⁷

Līdzās juridiski saistošam tiesiskajam regulējumam, ir pieņemtas arī MI ētikas vadlīnijas, kas nosaka ētikas pamatprincipus, kā arī sniedz ieteikumus kā nodrošināt ētisku un atbildīgu MI izmantošanu, bet nav juridiski saistošas. UNESCO Ieteikums par mākslīgā intelekta ētiku⁸⁸, ko pieņēma 2021. gadā, nosaka vairākus ētikas pamatprincipus: proporcionalitāte un kaitējuma nepieļaušana, drošība, taisnīgums un nediskriminācija, ilgtspēja, tiesības uz privātumu un datu aizsardzība, cilvēku uzraudzība, pārskatāmība un izskaidrojamība, atbildība. Taisnīguma un nediskriminācijas princips (28. – 30. punkts) nosaka, ka MI dalībniekiem jāpieliek visas saprātīgās pūles, lai līdz minimumam samazinātu un izvairītos no diskriminējošu vai neobjektīvu lietojumu un rezultātu pastiprināšanas vai saglabāšanas visā MI sistēmas dzīves ciklā, lai nodrošinātu šādu sistēmu taisnīgumu. Jābūt pieejamiem efektīviem tiesiskās aizsardzības līdzekļiem pret diskrimināciju un neobjektīvu algoritmisku noteikšanu (29. punkts). UNESCO Ieteikumā par mākslīgā intelekta ētiku ir noteikts arī pārredzamības un izskaidrojamības princips (37. – 41. punkts). Cita starpā tas paredz, ka personām ir jādara zināms, kad tās komunicē ar MI (38. punkts).

Attiecībā uz MI izmantošanu tiesu darbā, UNESCO Ieteikums par mākslīgā intelekta ētiku paredz, ka dalībvalstīm būtu jāuzlabo tiesu iestāžu spēja pieņemt lēmumus saistībā ar MI sistēmām saskaņā ar tiesiskumu un starptautiskajām tiesībām un standartiem,

⁸⁷ [Council of Europe. \(2025\). European Commission for the Efficiency of Justice \(CEPEJ\) working group on cyberjustice and artificial intelligence \(cepej-gt-cyberjust\) advisory board on artificial intelligence \(CEPEJ-AIAB\)](#)

⁸⁸ [UNESCO. \(2021\). Recommendation on the Ethics of Artificial Intelligence.](#)

tostarp attiecībā uz MI sistēmu izmantošanu savās apspriedēs, vienlaikus nodrošinot, ka tiek ievērots cilvēka uzraudzības princips. Ja tiesu iestādes izmanto MI sistēmas, ir nepieciešami pietiekami drošības pasākumus, lai cita starpā garantētu cilvēka pamattiesību, tiesiskuma, tiesu neatkarības, kā arī cilvēka uzraudzības principa aizsardzību un nodrošinātu uzticamu, uz sabiedrības interesēm un uz cilvēku orientētu MI sistēmu izstrādi un izmantošanu tiesu iestādēs (63. punkts).

Arī tiesas ir sākušas izstrādāt vadlīnijas MI sistēmu izmantošanai. Piemēram, Eiropas Savienības Tiesa 2023. gadā publicēja Mākslīgā intelekta stratēģiju.⁸⁹ Tā paredz, ka tad, kad ir ieviesti MI risinājumi, procedūras, metodes un pārvaldība, darbinieku informētībai un zināšanu līmenim ir jānodrošina, ka tiek ievēroti šādi principi:

- 1) taisnīgums, neitralitāte un nediskriminācija; 2) pārredzamība; 3) izsekojamība;
- 4) privātums un datu aizsardzība; 5) cilvēka uzraudzība; 6) nepārtraukta uzlabošana.

Taisnīguma, neitralitātes un nediskriminācijas princips nosaka, ka gan datiem, gan izveidotajiem vai pieņemtajiem algoritmiem jāizvairās no aizspriedumiem un jāvadās pēc taisnīguma un neitralitātes principa, lai visas puses tiesas vai administratīvā procesa laikā saņemtu vienlīdzīgu attieksmi. Izveidotajam vai izmantotajam MI risinājumam nevajadzētu diskriminēt nevienu indivīdu vai grupu tādu faktoru dēļ kā rase, dzimums vai sociālekonomiskais statuss.

Eiropas Padomes Eiropas Tiesiskuma efektivitātes komisija (CEPEJ) jau 2018. gadā publicēja Eiropas ētikas hartu par mākslīgā intelekta izmantošanu tiesu sistēmās un to vidē.⁹⁰ Tā nosaka piecus principus: 1) pamattiesību ievērošanas princips, kas paredz nodrošināt, lai MI rīku un pakalpojumu izstrāde un ieviešana atbilstu pamattiesībām; 2) nediskriminācijas principu, kas jo īpaši paredz novērst jebkādas diskriminācijas attīstību vai pastiprināšanos starp indivīdiem vai indivīdu grupām; 3) kvalitātes un drošības princips, kas paredz attiecībā uz tiesu lēmumu un datu apstrādi izmantot sertificētus avotus un nemateriālus datus ar modeļiem, kas izstrādāti daudznazaru veidā, drošā tehnoloģiskā vidē; 4) pārredzamības, neitralitātes un taisnīguma princips, kas uzliek pienākumu padarīt datu apstrādes metodes pieejamas un saprotamas, atļaut

⁸⁹ [Court of Justice of the European Union. \(2023\). Artificial Intelligence Strategy.](#)

⁹⁰ [European Commission for the Efficiency of Justice \(CEPEJ\) \(2018\). European ethical Charter on the use of Artificial Intelligence in judicial systems and their environment.](#)

ārējās revīzijas; 5) princips “lietotāja kontrolē”, kas paredz nepieļaut normatīvu pieeju un nodrošināt, ka lietotāji ir informēti dalībnieki un kontrolē izdarītās izvēles.

MI sistēmas, jo īpaši tās, kuru pamatā ir mašīnmācīšanās un dabiskās valodas apstrāde, kļūst arvien svarīgākas tiesās, un ģeneratīvā MI izmantošana tieslietu nozarē ievērojami pieaug.⁹¹ MI sistēmas var prognozēt juridisko strīdu iznākumu, sniegtot vērtīgu ieskatu juristiem un palīdzot tiesnešiem lēmumu pieņemšanā. Tomēr pašreizējām MI sistēmām ir būtiski ierobežojumi un tām nepieciešama cilvēka uzraudzība. Nav pilnībā autonomu MI sistēmu, kas spētu neatkarīgi darboties tiesās.⁹² Turpmāk ir aplūkoti daži piemēri no starptautiskās prakses, kas parāda, kā MI sistēmu izmantošana tiesvedības procesos var radīt diskriminācijas risku.

Prakses piemēri

COMPAS sistēma, lai novērtētu risku, ka persona izdarīs noziedzīgu nodarījumu

Neobjektivitāte attiecas uz sistemātisku novirzi no patiesās vērtības vai objektīva standarta, kā rezultātā rodas pastāvīga labvēlība vai aizspriedumi pret noteiktiem rezultātiem. Sistēmām, kas balstās uz datiem, ir nepieciešams apkopot un izmantot lielu datu kopumu. Tomēr sliktas kvalitātes, neprecīzi un/vai slikti pārrakstīti dati arī ietekmē rezultātus. Turklāt pašā sistēmā var būt iebūvētas kļūdas un/vai neobjektivitātes. Rezultātu var ietekmēt arī sistēmas izveidotāju subjektīvie viedokļi par to, kuri dati tiek izmantoti, iekļauti vai izslēgti, kā arī par to, kā tiem piešķirt nozīmi un kurai informācijai piešķirt prioritāti vai to samazināt līdz minimumam.

Viens no spilgtiem piemēriem, kas parāda, kā MI sistēmu izmantošana var radīt šāda veida aizspriedumus, ir ASV tiesu iestāžu izmantotā sistēma COMPAS - Koriģējošā likumpārkāpēju pārvaldības profilēšana alternatīvām sankcijām (Correctional Offender Management Profiling for Alternative Sanction – angļu val.), kas izrādījās pretrunīga iespējama recidīvisma novērtēšanā.⁹³ Šis rīks tika izmantots, lai piesprieztu personai sešu gadu cietumsodu par

⁹¹ [Council of Europe. \(2025\). European Commission for the Efficiency of Justice \(CEPEJ\) working group on cyberjustice and artificial intelligence \(cepej-gt-cyberjust\) advisory board on artificial intelligence \(CEPEJ-AIAB\)](#)

⁹² Turpat.

⁹³ [Golobardes, M. A., Grasso G., Iamiceli, P. Et. al. \(2024\). JuLIA Handbook. Artificial Intelligence, Judicial Decision-Making and Fundamental Rights.](#); [Orakhelashvili, A. \(2024\). Real World Cases of Annex III AI Act applications that posed risks to fundamental rights.](#)

automašīnas vadīšanu, kas izmantota apšaudē. Spriedums tika pārsūdzēts, norādot uz tiesību uz taisnīgu tiesu pārkāpumu. Apelācijas sūdzība tika pamatota ar to, ka sistēmas piešķirto vērtējumu nebija iespējams izvērtēt, un COMPAS sistēmas noteiktais sods pārkāpj personas tiesības uz individualizētu sodu, kas balstīts uz precīzu informāciju. COMPAS sistēmas veiktais riska novērtējums ir balstīts uz datiem, kas iegūti no apsūdzētā krimināllietas un no apsūdzētā pratināšanas, un paredz pirmstiesas recidīvisma, vispārēja recidīvisma un vardarbīga recidīvisma risku. Persona visos trijos rādītājos ieguva augstu vērtējumu. Taču, tā kā tiesnesis norādīja, ka izmantojis COMPAS ziņojumu, persona pieprasīja nozīmēt jaunu tiesas sēdi. Šis pieprasījums tika noraidīts, taču Viskonsinas Augstākās tiesa ieteica tiesnešiem, kuriem tika iesniegti COMPAS vērtējumi, saprast sistēmas vājās vietas. Tā kā tiesnesis lietā savu lēmumu pamatoja arī ar personas kriminālo pagātņi, jauna tiesas sēde tika atteikta.

COMPAS programmatūra tika izstrādāta tā, lai neņemtu vērā personu **rasi un etnisko piederību**, tāpēc bija sagaidāms, ka šāda veida dati neietekmēs recidīvisma riska novērtējumu. Tomēr nevalstiskās organizācijas ProPublica pētījums atklāja, ka etniskajai piederībai netieši bija nozīme un tā pat atsvēra citus skaidri iekļautus faktorus, ņemot vērā savstarpējas atsauces uz citiem datiem, piemēram, dzīvesvietu vai profesiju, kā arī tāpēc, ka noteiktas etniskās grupas bija pārstāvētas datos, kas izmantoti sistēmas apmācībai. Piemēram, afroamerikāņiem divu gadu laikā pēc sprieduma pasludināšanas tika piešķirts augsta recidīvisma risks, kas bija divreiz lielāks nekā citām etniskām grupām, tādējādi samazinot viņu atbrīvošanas iespējamību no apcietinājuma. Ir bijuši gadījumi, kad ieslodzītajiem ar praktiski nevainojamu ieslodzījuma vēsturi ir atteikta nosacīta atbrīvošana neprecīza COMPAS rādītāja dēļ, atstājot viņiem maz iespēju apstrīdēt lēmumu vai pat uzzināt, kāpēc tika pieņemts attiecīgais lēmums. Lai gan izstrādātāji to necentās panākt, algoritms uzskatīja citas etniskās grupas par tādām, kurām ir mazāka iespēja atkārtoti izdarīt nodarījumu. Tas parāda, ka aizspriedumu risks ir ievērojams un to nav viegli novērst. Mī rīki, kas apmācīti ar vēsturiskiem datiem, neizbēgami rada modeļus, kas atkārtoti to, kas ir noticis pagātnē.⁹⁴

⁹⁴ [Golobardes, M. A., Grasso G., Iamiceli, P. Et. al. \(2024\). JuLIA Handbook. Artificial Intelligence, Judicial Decision-Making and Fundamental Rights.](#)

COMPAS izmantošanas lieta arī parāda, ka riska novērtējumu kvalitāti (un precizitāti) nosaka savākto datu kvalitāte. Tiesa jo īpaši norādīja uz to, ka rīks tika apmācīts ar visā valstī apkopotiem datiem, neveicot savstarpēju pārbaudi ar datiem par Viskonsinas iedzīvotājiem. Tā rezultātā, iespējams, netika pienācīgi atspoguļotas konkrētā štata, kurā rīks tika izmantots, specifiskās iezīmes, piemēram, demogrāfiskie dati un sociālais, ekonomiskais un juridiskais statuss. Tādēļ Viskonsinas Augstākā tiesa pieprasīja tiesnešiem un citiem lietas dalībniekiem, lietās, kurās tika izmantots COMPAS, paziņot par sistēmas ierobežojumiem.

MI sistēmas apkopo lielu datu apjomu un atrod biežas atbilstības vai modeļus, un pamatojoties uz to, tās sniedz gala rezultātu vai izvaddatus. Vairumā gadījumu nav iespējams zināt, kā tās nonāk pie gala rezultāta. Tāpēc pastāv risks, ka MI, sniedzot izvaddatus, ir izmantojis modeļus, kas reproducē aizliegtu diskrimināciju, piemēram, rases, etniskās izcelsmes, sociālekonomiskās izcelsmes, reliģijas vai politisko uzskatu dēļ.⁹⁵ Kā piemēru varētu minēt gadījumu, kad MI sistēma, atbildot uz tiesneša jautājumu par šķiršanās pabalsta apmēru, kas pienākas konkrētai personai, ņemtu vērā, ka vairumā gadījumu noteiktas etniskās izcelsmes personām piešķirtie pabalsti ir mazāki, un tieši sasaistītu to ar šo faktu, nevis ar to, ka vairumā gadījumu iesaistītās personas bija pāri ar zemiem ienākumiem. Norāde par nelielu pabalsta apmēru, kas piešķirams lietā iesaistītajai personai, būtu diskriminējoša. MI sistēmu izmantošana var arī samazināt cilvēku diskriminācijas risku, piemēram, kas balstīta uz atsevišķa tiesneša aizspriedumiem. Tomēr tā nevar izvairīties no diskriminācijas, kas ir ietverta izmantotajos datos (ko bieži vien atspoguļo precedenti tiesu praksē), jo MI tiesu sistēmas nepieļauj esošo datu evolūciju, bet gan balsta savus rezultātus uz tiem un pat atkārto savus risinājumus.⁹⁶

⁹⁵ [European Union Agency for Fundamental Rights \(FRA\). \(2022\). Bias in Algorithms. Artificial Intelligence and Discrimination.](#)

⁹⁶ [Golobardes, M. A., Grasso G., Iamiceli, P. Et. al. \(2024\). JuLIA Handbook. Artificial Intelligence, Judicial Decision-Making and Fundamental Rights.](#)

Austrālijas Split-up sistēma ģimenes tiesību tiesās

MI ekspertu un juristu grupa ir izstrādājusi Split-Up sistēmu, ko izmanto Austrālijas ģimenes tiesību tiesās, lai paredzētu īpašuma strīdu iznākumu laulības šķiršanas un citos ģimenes tiesību jautājumos.⁹⁷

Split-Up sistēmu tiesneši izmanto, lai atbalstītu lēmumu pieņemšanu, palīdzot viņiem noteikt izlīgumā norādāmo īpašumu summu. Proti, sistēma palīdz tiesnesim noteikt, cik lielu procentuālo daļu no kopējā īpašuma katrai pusei vajadzētu saņemt, pamatojoties uz tādiem faktoriem kā iemaksas, ienākumu avoti un nākotnes vajadzības. Sistēma analizē 94 galvenos elementus, izmantojot statistikas metodes, kuru pamatā ir neironu tīkla arhitektūra. Pēc tam tiesnesis var ierosināt pieņemt galīgo īpašuma rīkojumu, pamatojoties uz algoritma veikto analīzi. Sistēmas mērķis ir arī sniegt skaidru pamatojumu lēmumiem.

Viena no problēmām saistībā ar neobjektivitāti, izmantojot tādas sistēmas kā Split-Up, ir tā, ka šajā kontekstā izmantotos datus (šķiršanās strīdiem parasti raksturīga **dzimumu nelīdztiesība**, un vēsturiskie dati var atspoguļot diskriminācijas modeli) MI sistēma var uztvert kā absolūtu patiesību. Tiesu iestāžu darbiniekiem jāapzinās šīs problēmas un riski, kas saistīti ar šāda veida MI sistēmām.⁹⁸

Singapūras tiesas transkripcijas sistēma

Singapūras tiesās tika ieviesta tiesas transkripcijas sistēma (iCTS), lai palielinātu tiesu efektivitāti, transkribējot tiesas sēdes reāllaikā un ļaujot tiesnešiem un pusēm nekavējoties pārskatīt mutvārdu liecības tiesā. Tas tiek panākts, izmantojot neironu tīklus, kas apmācīti ar valodu modeļiem un konkrētai jomai raksturīgiem terminiem (piemēram, juridisko terminoloģiju).⁹⁹ Jāatzīmē, ka runas atpazīšanas sistēmām var nedarboties labi, ja tās tiek pakļautas noteiktiem valodas akcentiem, kas noteiktos apstākļos var būt diskriminējoši. Tiesu iestāžu darbiniekiem ir jāapzinās šie trūkumi.¹⁰⁰

⁹⁷ [Stankovich, M, Feldfeber, I.; Quiroga Y. \(2023\). Global Toolkit on AI and the Rule of Law for the Judiciary.](#)

⁹⁸ Turpat.

⁹⁹ [Stankovich, M, Feldfeber, I.; Quiroga Y. \(2023\). Global Toolkit on AI and the Rule of Law for the Judiciary.](#); [Ng J. \(2020\). Legal Tech-ing Our Way to Justice.](#)

¹⁰⁰ [Stankovich, M, Feldfeber, I.; Quiroga Y. \(2023\). Global Toolkit on AI and the Rule of Law for the Judiciary.](#)

MI sistēmas jāprojektē tā, lai nodrošinātu taisnīgumu un izvairītos no aizspriedumiem, kas var novest pie diskriminējošiem rezultātiem. Ir ļoti svarīgi novērst aizspriedumus apmācības datos, algoritmos un lēmumu pieņemšanas procesos, lai novērstu netaisnīgu attieksmi pret noteiktām personām vai grupām vai to marginalizāciju.¹⁰¹ Savukārt tiesnesim jebkurā gadījumā ir jāpārbauda, vai nav pieļauta diskriminācija, ja tiek izmantota MI sistēma. MI sistēmas var būt vērtīgs palīgs tiesnešiem (pat palīdzot viņiem pārvarēt aizspriedumus), taču tās nevar ne aizstāt pašus tiesnešus, ne pilnībā aizstāt viņu darba tradicionālos rīkus.¹⁰²

¹⁰¹ Turpat.

¹⁰² Turpat.

7. Kopsavilkums

Mākslīgā intelektā akts klasificē MI sistēmas kā augsta riska, ņemot vērā, ka tās var radīt būtisku kaitējumu personu pamattiesībām, demokrātijai un tiesiskumam. Pētījumā analizētās četras augsta riska jomas un prakses piemēri atklāj vairākus problēmjautājumus, kas var traucēt pamattiesību, īpaši diskriminācijas novēršanas principa, ievērošanu.

Tiesībaizsardzības jomā saskaņā ar Mākslīgā intelekta aktu augstu risku rada MI sistēmas, ko paredzēts izmantot piecos veidos: 1) lai novērtētu risku, ka fiziska persona varētu kļūt par cietušo personu noziedzīgos nodarījumos; 2) kas atbalsta tiesībaizsardzības iestādes, piemēram, melu detektoru un līdzīgi rīki; 3) pierādījumu uzticamības izvērtēšanai noziedzīgu nodarījumu izmeklēšanas vai kriminālvajāšanas gaitā; 4) lai novērtētu risku, ka fiziska persona varētu izdarīt vai atkārtoti izdarīt nodarījumu, pamatojoties ne tikai uz fizisku personu profilēšanu, kas minēta Policijas direktīvas 3. panta 4. punktā, vai lai novērtētu fizisku personu vai grupu personības iezīmes un rakstura īpašības vai agrāku noziedzīgu rīcību; 5) Policijas direktīvas 3. panta 4. punktā minētajai fizisku personu profilēšanai noziedzīgu nodarījumu atklāšanas, izmeklēšanas vai kriminālvajāšanas gaitā. Mākslīgā intelekta akta III pielikumā noteiktas arī cita veida augsta riska MI sistēmu, ko izmanto tiesībaizsardzības iestādes: biometrija; migrācijas, patvēruma un robežkontroles pārvaldība; kā arī tiesvedība.

MI sistēmas, kas rada augstu risku tiesībaizsardzības jomā, ir svarīgi nošķirt no vairākiem aizliegtās MI prakses veidiem: reāllaika biometriskās tālidentifikācijas sistēmu izmantošanas publiski piekļūstamās vietās tiesībaizsardzības nolūkos; noziedzīgu nodarījumu riska novērtēšanas un prognozēšanas, kas balstīta vienīgi uz personas profilēšanu vai personisko īpašību un rakstura iezīmeju novērtēšanu. Aizliegums izmantot MI sistēmas noziedzīgu nodarījumu riska novērtēšanā un prognozēšanā jeb prognozējošajā policijas darbībā neattiecas uz MI sistēmām, kuras izmanto, lai atbalstītu cilvēka veiktu novērtējumu par personas iesaistīšanos noziedzīgā darbībā, kurš jau ir balstīts uz objektīviem un pārbaudāmiem faktiem. Proti, izšķiroša nozīme ir cilvēka virsvadības prasībai.

Mākslīgā intelekta aktā ir paredzēti vairāki vispārīgi izņēmumi no tā darbības jomas, kas ir būtiski, lai izprastu III pielikumā uzskaitīto augsta riska jomu praktisko piemērošanu, īpaši

noteikums, ka tas neattiecas uz MI sistēmām, ja tās tiek izmantotas vienīgi militāros, aizsardzības vai valsts drošības nolūkos.

Pastāv daudzi prakses piemēri, kas atklāj, kā dažāda veida MI sistēmas tiesībaizsardzības jomā var radīt diskriminācijas risku, galvenokārt **tautības, etniskās vai rases piederības** dēļ, kā arī pārredzamības un uzraudzības problēmas. Eiropā ir vairāki piemēri, kad tiek izmantotas prognozējošās policijas sistēmas, kas iespējams rada diskriminācijas risku. MI sistēmu izmantošanai tiesībaizsardzības jomā, tai skaitā prognozējošā policijas darbībā, ir īpaši raksturīgs neobjektivitātes risks, kas var rasties jebkurā MI sistēmas dzīves cikla posmā. Riski rodas no aizspriedumiem, kas iestrādāti MI sistēmu projektēšanas, izstrādes un ieviešanas gaitā, kas var veicināt diskrimināciju, pastiprināt sabiedrības nevienlīdzību un apdraudēt tiesībaizsardzības iestāžu darbību integritāti.

Kritiskā infrastruktūra ir noteikta kā viena no augsta riska jomām. Proti, saskaņā ar Mākslīgā intelekta aktu kā augsta riska tiek klasificētas MI sistēmas, ko paredzēts izmantot kā drošības sastāvdaļas kritiskās digitālās infrastruktūras pārvaldībā un darbībā, ceļu satiksmē vai ūdens, gāzes vai elektroapgādē vai apkures nodrošināšanā. MI sistēmas, ko izmanto kritiskās infrastruktūras jomā, var apdraudēt personu veselību un drošību, un tās pamatā nav saistītas ar diskriminācijas risku. MI sistēmu izmantošana kritiskajā infrastruktūrā ir saistīta ar citu augsta riska jomu – piekļuvi privātiem pamatpakalpojumiem un sabiedriskajiem pamatpakalpojumiem un pabalstiem un to izmantošanu. Piemēram, MI izmantošana, lai pieņemtu lēmumus, kas saistīti ar enerģētiku, rada bažas par algoritmisko neobjektivitāti, datu privātumu un atbildību. MI sistēmas, kas apmācītas, izmantojot vēsturiskus datus, kas atspoguļo sabiedrības nevienlīdzību, var novest pie diskriminējošiem rezultātiem resursu sadalē vai cenu noteikšanā.

MI sistēmu izmantošana rada bažas par kiberdrošības riskiem MI vadītā kritiskās infrastruktūras aizsardzībā un paplašina uzbrukuma iespējas kritiskajai infrastruktūrai. Kritiskās infrastruktūras īpašniekiem un valdītājiem ir jāzina, kā, kur un kāpēc MI sistēmas tiks izmantotas, lai novērtētu kontekstam un nozarei specifiskus riskus un iespējamo ietekmi uz drošību un aizsardzību.

Migrācijas, patvēruma un robežkontroles pārvaldības jomā augstu risku rada MI sistēmas, ko paredzēts izmantot četros veidos: 1) kā melu detektorus vai līdzīgus rīkus; 2) lai novērtētu risku, tostarp drošības risku, neatbilstīgas migrācijas risku vai veselības risku, ko rada fiziska persona, kas vēlas ieceļot vai ir ieceļojusi dalībvalsts teritorijā; 3) lai palīdzētu kompetentām publiskajām iestādēm izskatīt patvēruma, vīzu vai uzturēšanās atļauju pieteikumus, kā arī saistītās sūdzības, nosakot, vai fiziskās personas, kuras iesniegušas pieteikumu uz kādu no minētajiem statusiem, atbilst kritērijiem, tostarp veicot saistītu pierādījumu patiesuma novērtēšanu; 4) saistībā ar migrācijas, patvēruma vai robežkontroles pārvaldību nolūkā atklāt, atpazīt vai identificēt fiziskas personas, izņemot ceļošanas dokumentu verifikāciju.

MI sistēmu precizitāte, nediskriminējošais raksturs un pārredzamība šajā jomā ir īpaši svarīga, lai tiktu ievērotas personu pamattiesības. Tās nedrīkst lietot, lai apietu starptautiskās saistības, pārkāptu neizraidīšanas principu vai liegtu drošas un likumīgas iespējas iekļūt ES teritorijā, tostarp tiesības uz starptautisko aizsardzību.

Migrācijas, patvēruma un robežkontroles pārvaldības jomā ir nepieciešams nošķirt augsta riska MI sistēmu izmantošanu no aizliegtas MI prakses. Augsta riska MI sistēmas ir jānošķir no aizliegtas prakses, proti, kad MI sistēmu lieto fizisku personu radītā riska novērtējuma veikšanai, kura mērķis ir novērtēt vai prognozēt personas noziedzīga nodarījuma izdarīšanas risku, balstoties vienīgi uz personas profilēšanu. Šāda prakse pārkāpj nevainīguma prezumpcijas principu, saskaņā ar kuru personas ir jāvērtē pēc to faktiskās uzvedības, nevis pēc MI sistēmu prognozētas uzvedības, pamatojoties tikai uz to profilēšanu, personiskajām īpašībām vai rakstura iezīmēm. Aizliegums ir būtisks Vīzu kodeksa kontekstā, kur trešās valsts valstspiederīgajam var atteikt ieceļošanu vai vīzas piešķiršanu, ja šī persona tiek uzskatīta par draudu sabiedriskajai kārtībai vai iekšējai drošībai.

Migrācijas, patvēruma un robežkontroles pārvaldības jomā ir svarīgi arī citi MI prakses aizliegumi: nemērķēta sejas attēlu vākšana no interneta vai videomateriāla, kas iegūts ar videonovērošanas sistēmas palīdzību, nolūkā izveidot vai paplašināt datubāzes; fizisku personu biometriskā kategorizācija, kas ļauj noteikt vai izsecināt viņu rasi, politiskos uzskatus, dalību arodbiedrībās, reliģiskos vai filozofiskos uzskatus, dzimumdzīvi vai seksuālo orientāciju; reāllaika biometriskās tālidentifikācijas izmantošana.

Pastāv daudzi ārvalstu prakses piemēri, kas parāda, kā MI sistēmu izmantošana migrācijas, patvēruma un robežkontroles pārvaldības jomā var pārkāpt diskriminācijas aizlieguma principu. Spilgts piemērs, kas parāda šādu sistēmu radīto risku ir iBorderCtrl automatizētā melu noteikšanas sistēma imigrācijas kontrolei. Šāda veida sistēmas, kurām trūkst zinātniska pamatojuma, būtiski apdraud cilvēktiesības, tai skaitā diskriminācijas aizlieguma principa ievērošanu.

Gan Eiropā, gan citās valstīs MI sistēmas tiek izmantotas, lai palīdzētu izskatīt patvēruma, vīzu vai uzturēšanās atļauju pieteikumus, kas rada diskriminācijas riskus. MI sistēmas tiek arvien vairāk izmantotas valodas analīzei, lai noteiktu patvēruma meklētāju izcelsmes valsti. Lai gan MI sistēmas pēdējās desmitgades laikā ir ievērojami attīstījušās, tās joprojām var būt neprecīzas un neobjektīvas. Neprecīzi un/vai neobjektīvi MI novērtējumi vai ieteikumi par svarīgiem patvēruma pieteikumu aspektiem var izraisīt būtiskus pamattiesību pārkāpumus. Neobjektīvi algoritmi var izraisīt diskrimināciju pret **sievietēm, etniskajām minoritātēm, vecāka gadagājuma cilvēkiem un citām grupām**. Tā kā MI sistēmas tiek izstrādātas, izmantojot liela apjoma datus, šo datu kvalitāte ir svarīga, lai garantētu šo sistēmu precizitāti un uzticamību. Kļūdas datu analīzē vai interpretācijā var izraisīt nepareizus secinājumus par pieteikuma iesniedzēja izcelsmes valsti, kas var novest pie negodīgiem lēmumiem, kas var radīt bīstamas sekas attiecīgās personas dzīvībai un drošībai.

Tiesvedības un demokrātijas procesu jomā augstu risku rada MI sistēmas, ko paredzēts izmantot divos veidos: 1) lai palīdzētu tiesu iestādei izpētīt un interpretēt faktus vai tiesību normas un piemērot tiesību normas konkrētam faktu kopumam, vai izmantot līdzīgā veidā strīdu alternatīvā izšķiršanā; 2) lai ietekmētu vēlēšanu vai referenduma rezultātus vai fizisku personu uzvedību balsošanā, kad tās balso vēlēšanās vai referendumos. Šīs MI sistēmas ir klasificētas kā augsta riska sistēmas, ņemot vērā to potenciāli ievērojamo ietekmi uz demokrātiju, tiesiskumu, individuālajām brīvībām, kā arī tiesībām uz efektīvu tiesību aizsardzību un taisnīgu tiesu.

Ir jānošķir augsta riska MI sistēmas, ko izmanto tiesvedības un demokrātijas procesu jomā, no aizliegtas MI prakses, kas var apdraudēt demokrātijas procesus: kaitējošas manipulācijas, maldināšanas, sublimināliem paņēmieniem; cilvēku neaizsargātības izmantošanas. Mākslīgā intelekta aktā noteiktie aizliegumi atbilst arī ES datu aizsardzības tiesību aktiem, kuru mērķis

ir aizsargāt datu subjektu personas datus un pamattiesības. Līdzīgi Mākslīgā intelekta akts neietekmē praksi, kas ir aizliegta ar ES diskriminācijas novēršanas tiesību aktiem.

Mākslīgā intelekta akts ir vērsts uz pieaugošām bažām par dziļviltajumiem. Lai gan tas klasificē MI sistēmas, kas rada dziļviltajumus, kā ierobežota riska sistēmas, paredzot pārredzamības pienākumu, tajā pašā laikā MI sistēmas, ko izmanto, lai ietekmētu vēlēšanas, referendumu vai fizisku personu balsošanas uzvedību, tiek uzskatītas par augsta riska sistēmām, kā arī to izmantošana var būt aizliegta MI prakse.

MI, tostarp ģeneratīvais MI, var atstāt būtisku ietekmi uz sabiedrību un demokrātiju kopumā, jo tas palielina informācijas manipulācijas, dezinformācijas un viltus ziņu risku. Lai cīnītos ar šiem draudiem ne tikai ES, bet arī nacionālajā līmenī tiek pilnveidots tiesiskais regulējums. Mākslīgā intelekta akts papildina citus ES tiesību aktus: Digitālo pakalpojumu aktu un Regulu (ES) 2024/900 par politiskās reklāmas pārredzamību un mērķorientēšanu. Latvijā, lai regulētu MI izmantošanu priekšvēlēšanu aģitācijā, 2024. gadā Saeima pieņēma grozījumus Priekšvēlēšanu aģitācijas likumā¹⁰³, kas paredz pienākumu informēt vēlētājus par MI izmantošanu aģitācijas materiālu sagatavošanā. Lai cīnītos ar dziļviltajumiem vēlēšanu procesā, tika izdarīti grozījumi Krimināllikumā.

MI sistēmas, ko paredzēts izmantot tiesu darbā, tiek klasificētas kā augsta riska sistēmas, jo īpaši, lai mazinātu potenciālu neobjektivitātes, kļūdu un nepārredzamības risku. MI rīku lietošana var atbalstīt tiesnešu lēmumu pieņemšanas pilnvaras vai tiesu neatkarību, bet tai nevajadzētu to aizstāt. Proti, galīgajai lēmumu pieņemšanai arī turpmāk jābūt cilvēka veiktai darbībai

2024. gadā pieņemtā EP Konvencija paredz vispārīgu pienākumu aizsargāt pamattiesības un ievērot demokrātisko procesu integritāti un tiesiskumu. Mākslīgā intelekta aktam un EP Konvencijai ir kopīgi mērķi aizsargāt pamattiesības un veicināt atbildīgu MI izstrādi, bet tie atšķiras pēc to juridiskā rakstura, darbības jomas un izpildes mehānismiem. Lai gan Mākslīgā intelekta akts paredz svarīgus soļus augsta riska MI sistēmu regulēšanā, tā

¹⁰³ Grozījumi Priekšvēlēšanu aģitācijas likumā. Pieņemts 24.10.2024. Latvijas Vēstnesis, 06.11.2024, Nr. 217.

pieeja tiesu MI sistēmām ir jāpilnveido. Tas rada bažas par atbilstošiem tiesās izmantoto MI sistēmu uzraudzības mehānismiem.

Līdzās juridiski saistošam tiesiskajam regulējumam, ir pieņemtas arī MI ētikas vadlīnijas, kas nosaka ētikas pamatprincipus, kā arī sniedz ieteikumus kā nodrošināt ētisku un atbildīgu MI izmantošanu, tai skaitā tiesvedības un demokrātisko procesu jomā. UNESCO ieteikums par mākslīgā intelekta ētiku, attiecībā uz MI izmantošanu tiesu darbā uzsver nepieciešamību nodrošināt cilvēka pamattiesību, tiesiskuma, tiesu neatkarības, kā arī cilvēka uzraudzības principa aizsardzību. Arī tiesas ir sākušas izstrādāt vadlīnijas MI sistēmu izmantošanai tās darbā, piemēram, Eiropas Savienības Tiesa 2023. gadā publicēja Mākslīgā intelekta stratēģiju. Eiropas Padomes Eiropas Tiesiskuma efektivitātes komisija (CEPEJ) 2018. gadā publicēja Eiropas ētikas hartu par MI izmantošanu tiesu sistēmās un to vidē.

MI sistēmas, kuru pamatā ir mašīnmācīšanās un dabiskās valodas apstrāde, kļūst arvien svarīgākas tiesās, un ģeneratīvā MI izmantošana tieslietu nozarē ievērojami pieaug. Tomēr pašreizējām MI sistēmām ir būtiski ierobežojumi un tām nepieciešama cilvēka uzraudzība.

MI sistēmas jāprojektē tā, lai nodrošinātu taisnīgumu un izvairītos no aizspriedumiem, kas var novest pie diskriminējošiem rezultātiem. Ir svarīgi novērst aizspriedumus apmācības datos, algoritmos un lēmumu pieņemšanas procesos, lai novērstu netaisnīgu attieksmi pret noteiktām personām vai grupām. Ja tiek izmantota MI sistēmas, tiesnešiem ir jāpārbauda, vai nav pieļauta diskriminācija. MI sistēmas var būt vērtīgs palīgs tiesnešiem, taču tās nevar aizstāt tiesnešus.

Rekomendācijas

Zemāk ir sniegti vairāki ieteikumi, lai nodrošinātu pamattiesību, īpaši diskriminācijas aizlieguma principa, ievērošanu un atbildīgu MI sistēmu izmantošanu augsta riska jomās.¹⁰⁴

¹⁰⁴ Daļa no ieteikumiem ir balstīta cita starpā uz 2024. gada Pētījumu un [Europol \(2025\), AI bias in law enforcement - A practical guide, Europol Innovation Lab observatory report, Publications Office of the European Union, Luxembourg](#).

- **Atbilstības nodrošināšana Mākslīgā intelekta aktā un citos tiesību aktos**

noteiktajām prasībām. Augsta riska MI sistēmas var radīt būtisku kaitējumu personu veselībai, drošībai un pamattiesībām. Ņemot vērā augsta riska MI sistēmu iespējamās kaitīgās sekas, tām ir jāatbilst stingrākām prasībām, kas ir noteiktas Mākslīgā intelekta aktā, tai skaitā, riska pārvaldība (9. pants), pārredzamība (13. pants), cilvēka virsvadība (14. pants), kiberdrošība, precizitāte un noturība (15. pants), datu kvalitāte un pārvaldība, sistēmu apmācība un testēšana (10. pants), reģistrācija ES datubāzē (49. un turpmākie panti) un iepriekšējs ietekmes uz pamattiesībām novērtējums (27. pants). Tiesībaizsardzības un citām valsts iestādēm ir svarīgi šīs prasības praktiski ieviest un konsekventi ievērot savā darbībā, īpaši izmantojot MI sistēmas augsta riska jomās.

Lai novērstu MI sistēmu izmantošanas rezultātā radītos riskus, būtiska nozīme ir arī citiem ES tiesību aktiem, īpaši datu aizsardzības regulējumam, kas ir piemērojams MI sistēmām, kas apstrādā personas datus.

- **Augsta riska un aizliegtas MI prakses nošķiršana.** Mākslīgā intelekta akts iedala MI sistēmas atkarībā no dažādiem to radītā riska līmeņiem, kurus ir svarīgi praksē pareizi nošķirt. Dažos gadījumos tādu MI sistēmu izmantošanu, kas klasificētas kā augsta riska sistēmas, var uzskatīt par aizliegtu MI praksi, ja ir izpildīti visi viena vai vairāku Mākslīgā intelekta akta 5. pantā paredzēto aizliegumu nosacījumi. MI sistēmas, kas saskaņā ar Mākslīgā intelekta akta III pielikumu ir norādītas kā augsta riska, var tikt klasificētas kā aizliegtas saskaņā ar Mākslīgā intelekta akta 5. pantu, kā arī citiem tiesību aktiem. Mākslīgā intelekta akta III pielikumā, uzskaitot atsevišķas augsta riska jomas (piemēram, tiesībaizsardzības, migrācijas, patvēruma un robežkontroles pārvaldības jomu), ir ietverta arī norāde, ka augstu risku rada attiecīgo MI sistēmu izmantošana “ciktāl to izmantošanu atļauj attiecīgie ES vai valsts tiesību akti”. Mākslīgā intelekta akts nesamazina aizsardzību, tostarp prakses aizliegumus, ko nosaka diskriminācijas novēršanas un datu aizsardzības tiesiskais regulējums, bet gan papildina esošo regulējumu. Būtu nepieciešams skaidras vadlīnijas, kas palīdzētu noteikt, vai MI sistēmu izmantošana ir uzskatāma par aizliegtu praksi vai arī rada augstu risku.

Eiropas Komisijai būtu jāizmanto Mākslīgā intelekta aktā paredzētās iespējas pārskatīt aizliegto MI sistēmu sarakstu, kā arī Mākslīgā intelekta akta III pielikumu, lai atjauninātu augsta riska MI sistēmu sarakstu atbilstoši tehnoloģiju attīstībai un sociālajām vajadzībām.

- **Datu kvalitāte.** Lai novērstu MI sistēmu izmantošanas radītos diskriminācijas riskus, ir jānodrošina atbilstība datu kvalitātes un pārvaldības prasībām. Mākslīgā intelekta akts nosaka, ka apmācības, validēšanas un testēšanas datu kopām ir jābūt atbilstošām, pietiekami reprezentatīvām un, cik vien iespējams, bez kļūdām un pilnīgām, ņemot vērā MI sistēmas paredzēto nolūku. Īpaša uzmanība ir jāpievērš tādas iespējamās neobjektivitātes mazināšanai datu kopās, kura varētu ietekmēt personu veselību un drošību, kurai varētu būt negatīva ietekme uz pamattiesībām vai kura varētu izraisīt diskrimināciju, kas ir aizliegta ES tiesību aktos, jo īpaši, ja izvaddati ietekmē ievaddatus turpmākām darbībām (atgriezeniskās saites cilpas). Pirms MI sistēmu ieviešanas visām pieejamajām datu kopām jāveic stingrs veikspējas un ietekmes novērtējums, kā arī neobjektivitātes testēšana. Strādājot ar iepriekš apmācītiem modeļiem, ne vienmēr var būt iespējams piekļūt sākotnējiem apmācības datiem vai tos pārbaudīt. Šādos gadījumos jāpaļaujas uz pieejamo dokumentāciju vai metadatiem un jāpārbauda modeļa rezultāti, lai noteiktu neobjektivitātes indikatorus. Ja tiek veikta iepriekš apmācīta modeļa precizēšana, ir rūpīgi jānovērtē un jāpārbauda precizējošo datu kopu, lai noteiktu neobjektivitāti pirms precizēšanas procesa un tā laikā, lai pārlicinātos, ka netiek ieviestas jaunas un pastiprinātas esošās neobjektivitātes.
- **Cilvēka virsvadība.** Lai nodrošinātu atbildīgu MI izmantošanu, ir jānodrošina cilvēka virsvadība, kas nozīmē, ka MI sistēmas darbojas tādā veidā, ko cilvēki var pienācīgi kontrolēt un uzraudzīt. Cilvēka virsvadība ir jāuztic personām, kurām ir nepieciešamā kompetence, apmācība un pilnvaras un kuras var pareizi izprast MI sistēmas spējas un ierobežojumi, interpretēt tās iznākumu un novērst automatizācijas neobjektivitātes risku. Personām būtu jāvar iejaukties, lai izvairītos no negatīvām sekām vai riskiem vai apturētu MI sistēmas izmantošanu, ja tā nedarbojas, kā paredzēts.

Cilvēka veiktai novērtēšanai ir jābūt neatņemamai MI sistēmas novērtēšanas sastāvdaļai. Taisnīguma pārbaude un pēcapstrādes neobjektivitātes mazināšanas metodes jāpiemēro gan MI sistēmas rezultātiem, gan galīgajiem lēmumiem, ko pieņem eksperti, kuri paļaujas uz šiem rezultātiem. Galu galā tās ir personas, kurām ir jānosaka, kā rīkoties, pamatojoties uz MI ģenerēto informāciju un ieteikumiem.

- **Pārredzamība.** Pārredzamība nozīmē, ka MI sistēmas tiek projektētas un lietotas tā, lai būtu iespējama pienācīga izsekojamība un izskaidrojamība un vienlaikus, lai cilvēki apzinātos, ka sazinās vai ir saskarē ar MI sistēmu, un lai pienācīgi informētu uzturētājus par MI sistēmas spējām un ierobežojumiem un skartās personas – par to tiesībām. Mākslīgā intelekta akts skaidro, ka augsta riska MI sistēmas ir jāprojektē tā, lai uzturētāji varētu saprast, kā MI sistēma darbojas, izvērtēt tās funkcionalitāti un izprast tās stiprās puses un ierobežojumus. Pārredzamībai, tostarp pievienotajai lietošanas instrukcijai, būtu jāpalīdz uzturētājiem lietot sistēmas un būtu jāpalīdz tiem pieņemt informētus lēmumus. Uzturētājiem cita starpā vajadzētu būt atbilstošām iespējām pareizi izvēlēties sistēmu, ko tie plāno lietot, ņemot vērā pienākumus, kas uz tiem attiecas, tiem vajadzētu būt izglītotiem par paredzētajiem un aizliegtajiem lietojumiem, un tiem būtu MI sistēma jālieto pareizi un atbilstoši. Lai uzlabotu lietošanas instrukcijā iekļautās informācijas skaidrību un piekļūstamību, attiecīgā gadījumā būtu jāiekļauj ilustratīvi piemēri, piemēram, par ierobežojumiem un par paredzētajiem un aizliegtajiem MI sistēmas lietojumiem. Nodrošinātājiem būtu jānodrošina, ka visa dokumentācija, tostarp lietošanas instrukcija, ietver jēgpilnu, visaptverošu, pieejamu un saprotamu informāciju, ņemot vērā uzturētāju vajadzības un paredzamās zināšanas. Lietošanas instrukcijām vajadzētu būt pieejamām uzturētājiem viegli saprotamā valodā.

Mākslīgā intelekta akts paredz arī tiesības saņemt skaidrojumu par individuālu lēmumu pieņemšanu. Jebkurai skartajai personai, uz kuru attiecas lēmums, ko uzturētājs pieņēmis, pamatojoties uz augsta riska MI sistēmas radīto iznākumu, un kas rada juridiskas sekas vai līdzīgā veidā būtiski ietekmē minēto personu tādā veidā, ka tā uzskata, ka tas nelabvēlīgi ietekmē tās veselību, drošību vai pamattiesības, ir tiesības saņemt no uzturētāja skaidrus un jēgpilnus paskaidrojumus par MI sistēmas

lomu lēmumu pieņemšanas procedūrā un pieņemtā lēmuma galvenajiem elementiem. Tas kalpo kā garantija tiesībām uz efektīvu tiesisko aizsardzību.

Pārredzamības pienākums ir attiecināms ne tikai attiecībā uz augsta riska sistēmām. Mākslīgā intelekta akts nosaka pārredzamības pienākumu, liekot dziļviltojumu veidotājiem informēt sabiedrību par sava darba mākslīgo raksturu. MI sistēmas, kura ģenerē vai manipulē attēla, audio vai video saturu, kas ir dziļviltojums, uzturētājiem ir skaidrā un nošķiramā veidā jāizpauž, ka saturs ir mākslīgi ģenerēts vai manipulēts. Līdzīgi arī tādas MI sistēmas, kura ģenerē vai manipulē tekstu, ko publicē nolūkā informēt sabiedrību par sabiedrības interešu jautājumiem, uzturētāji informē, ka teksts ir mākslīgi ģenerēts vai manipulēts. Šo pienākumu nepiemēro, ja lietošana ir likumiski atļauta noziedzīga nodarījuma atklāšanai, novēršanai, izmeklēšanai vai kriminālvajāšanai vai ja MI ģenerētu saturu ir pārskatījis cilvēks vai tam ir veikta redakcionāla kontrole un ja fiziska vai juridiska persona ir redakcionāli atbildīga par satura publikāciju.

- **Dokumentācija.** Lai nodrošinātu atbildību, izsekojamību un palīdzētu noteikt, kur var rasties neobjektivitāte, ir jāuztur detalizēta dokumentācija par visiem MI dzīves cikla posmiem, sākot no problēmas definēšanas līdz izstrādei un ieviešanai, tostarp par lēmumiem, kas pieņemti, pamatojoties uz MI sistēmas rezultātiem.
- **Nepārtraukta neobjektivitātes novērtēšana un mazināšana.** Jāievieš regulāra MI modeļu testēšana un atkārtota novērtēšana visā to dzīves ciklā, lai atklātu un mazinātu neobjektivitāti. Tas ietver lēmumu pieņemšanas procesu un to personu, kuras ietekmē MI izstrādi, rīcības analīzi. Lai nodrošinātu, ka novērtējumi ir visaptveroši un līdzsvaroti, ir jāveicina daudzu ieinteresēto personu iesaiste, kuru vidū būtu MI praktiķi un analītiķi, juridiskās un atbilstības komandas, datu zinātnieki un inženieri, MI un mašīnmācīšanās pētnieki, sociālie zinātnieki, ētikas speciālisti, datu aizsardzības iestādes, MI uzraudzības iestādes, politikas eksperti, pilsoniskā sabiedrība un interešu aizstāvības grupas. Tas ļauj nodrošināt plašāku perspektīvu un palīdz identificēt iespējamās nepilnības, padarot MI sistēmas taisnīgākas.

- **Regulāras apmācības un informētības veicināšana.** Mākslīgā intelekta akts uzliek pienākumu MI sistēmu nodrošinātājiem un uzturētājiem nodrošināt pietiekamu MI prasmes līmeni darbiniekiem un citām personām, kuras to vārdā iesaistītas MI sistēmu darbībā un izmantošanā. MI prasme ir prasmes, zināšanas un izpratne par Mākslīgā intelekta aktā noteiktajām tiesībām un pienākumiem, kas ļauj veikt MI sistēmu informēto uzturēšanu, kā arī izpratne par MI iespējām un riskiem un par potenciālo kaitējumu, ko tas var radīt.

Ir jāveic pastāvīgas apmācības visiem tiesībsardzības iestāžu un citu iestāžu darbiniekiem, kas iesaistīti MI rīku lietošanā, lai padziļinātu viņu izpratni par MI tehnoloģijām, neobjektivitātes sekām un taisnīguma rādītāju nozīmi. Šādās apmācībās jāuzsver cilvēka novērtējuma nozīme MI rezultātu pārskatīšanā, to atbildīga interpretācija un neobjektivitātes iespējamība. Lai tiesībsardzības un citas iestādes atbildīgi izmantotu MI sistēmas, ir svarīgi katram gadījumam veikt informēto un atsevišķu analīzi. Ir ieteicams regulāri apmācīt darbiniekus par dažādiem MI aizspriedumiem un to saistību ar taisnīguma rādītājiem un aizspriedumu mazināšanas metodēm.

Izmantotie avoti

Tiesību akti

Starptautiskie un Eiropas Savienības tiesību akti

Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law. 05.09.2024. <https://rm.coe.int/1680afae3c>

ANO Konvenciju par bēgļa statusu, kas grozīta ar 1967. gada 31. janvāra protokolu

Eiropas Savienības Pamattiesību harta. Pieņemta 07.12.2000. *OV C* 2020/239, 07.06.2016. http://data.europa.eu/eli/treaty/char_2016/oj

Eiropas Parlamenta un Padomes Regula (ES) 2024/900 (2024. gada 13. marts) par politiskās reklāmas pārredzamību un mērķorientēšanu, *OV L*, 2024/900, 20.3.2024. <http://data.europa.eu/eli/reg/2024/900/oj>

Eiropas Parlamenta un Padomes Regula (ES) 2024/1689 (2024. gada 13. jūnijs), ar ko nosaka saskaņotas normas mākslīgā intelekta jomā un groza Regulas (EK) Nr. 300/2008, (ES) Nr. 167/2013, (ES) Nr. 168/2013, (ES) 2018/858, (ES) 2018/1139 un (ES) 2019/2144 un Direktīvas 2014/90/ES, (ES) 2016/797 un (ES) 2020/1828 (Mākslīgā intelekta akts). *OV L* 2024/1689, 12.07.2024. <http://data.europa.eu/eli/reg/2024/1689/oj>

Komisijas deleģētā regula (ES) 2023/2450 (2023. gada 25. jūlijs), ar ko papildina Eiropas Parlamenta un Padomes Direktīvu (ES) 2022/2557, izveidojot pamatpakalpojumu sarakstu. *OV L*, 2023/2450, 30.10.2023. http://data.europa.eu/eli/reg_del/2023/2450/oj

Eiropas Parlamenta un Padomes Regula (ES) 2022/2065 (2022. gada 19. oktobris) par digitālo pakalpojumu vienoto tirgu un ar ko groza Direktīvu 2000/31/EK (Digitālo pakalpojumu akts, *OV L*, 227, 27.10.2022. <http://data.europa.eu/eli/reg/2022/2065/oj>

Eiropas Parlamenta un Padomes Regula (ES) 2016/679 (2016. gada 27. aprīlis) par fizisku personu aizsardzību attiecībā uz personas datu apstrādi un šādu datu brīvu apriti un ar ko atceļ Direktīvu 95/46/EK (Vispārīgā datu aizsardzības regula). *OV L* 119, 04.05.2016. <http://data.europa.eu/eli/reg/2016/679/oj>

Eiropas Parlamenta un Padomes Regula (EK) Nr. 810/2009 (2009. gada 13. jūlijs), ar ko izveido Kopienas Vīzu kodeksu (Vīzu kodekss), *OV L* 243, 15.9.2009. <http://data.europa.eu/eli/reg/2009/810/oj>

Eiropas Parlamenta un Padomes Direktīva (ES) 2022/2557 (2022. gada 14. decembris) par kritisko vienību noturību un Padomes Direktīvas 2008/114/EK atcelšanu. *OV L* 333, 27.12.2022. <http://data.europa.eu/eli/dir/2022/2557/oj>

Eiropas Parlamenta un Padomes Direktīva (ES) 2016/680 (2016. gada 27. aprīlis) par fizisku personu aizsardzību attiecībā uz personas datu apstrādi, ko veic kompetentās iestādes, lai novērstu, izmeklētu, atklātu noziedzīgus nodarījumus vai sauktu pie atbildības par tiem vai

izpildītu kriminālsodus, un par šādu datu brīvu apriti, ar ko atceļ Padomes Pamatlēmumu 2008/977/TI. OV L 119, 04.05.2016. <http://data.europa.eu/eli/dir/2016/680/oj>

Eiropas Parlamenta un Padomes Direktīva 2013/32/ES (2013. gada 26. jūnijs) par kopējām procedūrām starptautiskās aizsardzības statusa piešķiršanai un atņemšanai (pārstrādāta redakcija). OV L 180, 29.6.2013., <http://data.europa.eu/eli/dir/2013/32/oj>

UNESCO. (2021). Recommendation on the Ethics of Artificial Intelligence. <https://unesdoc.unesco.org/ark:/48223/pf0000381137>

Latvijas tiesību akti

Latvijas Republikas Satversme. Pieņemta 15.02.1922. (spēkā no 07.11.1922.). *Latvijas Vēstnesis*, 01.07.1993., Nr. 43.

Grozījumi Priekšvēlēšanu aģitācijas likumā. Pieņemts 24.10.2024. *Latvijas Vēstnesis*, 06.11.2024, Nr. 217.

Grozījumi Krimināllikumā. Pieņemts 26.09.2024. *Latvijas Vēstnesis*, 08.10.2024., Nr. 196.

Grozījumi Krimināllikumā. Pieņemts 09.05.2024. *Latvijas Vēstnesis*, 21.05.2024., Nr. 97.

Literatūra un citi materiāli

AIAAIC. (2022). US prison inmate AI call monitoring plan raises privacy, bias concerns. <https://www.aiaaic.org/aiaaic-repository/ai-algorithmic-and-automation-incidents/verus-prison-inmate-call-monitoring>

Algorithm Watch. (2020). How Dutch activists got an invasive fraud detection algorithm banned. <https://algorithmwatch.org/en/syri-netherlands-algorithm/>

Amnesty International. (2020). We sense trouble: Automated Discrimination and Mass Surveillance in Predictive Policing in the Netherlands. <https://www.amnesty.org/en/documents/eur35/2971/2020/en/>

Barkāne, I. (2024). Mākslīgais intelekts un diskriminācijas aspekti. Tiesībsarga birojs. <https://www.tiesibsargs.lv/resource/maksliga-intelekta-sistemas-un-diskriminācijas-aspekti/>

Barkane, I. (2023). Cilvēktiesību nozīme mākslīgā intelekta laikmetā. Privātums, Datu aizsardzība un regulējums masveida novērošanas novēršanai. Rīga: LU Akadēmiskais apgāds. <https://doi.org/10.22364/cnmil.23>

Brouwer, E. (2024). EU's AI Act and migration control. Shortcomings in safeguarding fundamental rights, *VerfBlog*. <https://dx.doi.org/10.59704/a4de76df20e0de5a>

Beduschi A. (2020). International migration management in the age of artificial intelligence, *Migration Studies*, 9(3), 576-596, <https://doi.org/10.1093/migration/mnaa003>

Council of Europe. (2025). European Commission for the Efficiency of Justice (CEPEJ) working group on cyberjustice and artificial intelligence (cepej-gt-cyberjust) advisory board on artificial intelligence (CEPEJ-AIAB) 1st AIAB Report on the use of artificial intelligence (ai) in the judiciary based on the information contained in the resource centre on cyberjustice and AI. <https://rm.coe.int/cepej-aiab-2024-4rev5-en-first-aiab-report-2788-0938-9324-v-1/1680b49def>

Court of Justice of the European Union. (2023). Artificial Intelligence Strategy. https://curia.europa.eu/jcms/upload/docs/application/pdf/2023-11/cjeu_ai_strategy.pdf

EDRi. (12 January, 2021). Re: Open letter: Civil society call for the introduction of red lines in the upcoming European Commission proposal on Artificial Intelligence. <https://edri.org/wp-content/uploads/2021/01/EDRi-open-letter-AI-red-lines.pdf>

Egbert, S., Krasmann, S. (2019). Predictive policing: Not yet, but soon preemptive? *Policing and Society*, 30(8), 905–919. <https://doi.org/10.1080/10439463.2019.1611821>

Eiropas Komisija. (2025). Komisijas paziņojums. Komisijas Pamatnostādnes par aizliegtu mākslīgā intelekta praksi, kas noteikta Regulā (ES) 2024/1690 (MI aktā). <https://digital-strategy.ec.europa.eu/lv/library/commission-publishes-guidelines-prohibited-artificial-intelligence-ai-practices-defined-ai-act>

Emian, Y. (2025.) Digital Borders and Racial Codes in AI Migration Control. <https://www.oba.org/digital-borders-and-racial-codes-in-ai-migration-control/>

European Commission for the Efficiency of Justice (CEPEJ) (2018). European ethical Charter on the use of Artificial Intelligence in judicial systems and their environment. <https://www.europarl.europa.eu/cmsdata/196205/COUNCIL%20OF%20EUROPE%20-%20European%20Ethical%20Charter%20on%20the%20use%20of%20AI%20in%20judicial%20systems.pdf>

European Parliament. (2025). Artificial intelligence in asylum procedures in the EU. [https://www.europarl.europa.eu/RegData/etudes/BRIE/2025/775861/EPRS_BRI\(2025\)775861_EN.pdf](https://www.europarl.europa.eu/RegData/etudes/BRIE/2025/775861/EPRS_BRI(2025)775861_EN.pdf)

European Union Agency for Asylum (EUAA). (2023). Strategy on Digital Innovation in Asylum Procedures and Reception Systems. https://euaa.europa.eu/sites/default/files/publications/2023-10/2023_EUAA-Strategy-on-Digital-Innovation-in-Asylum-Procedures-and-Reception-Systems_EN.pdf

European Union Agency for Asylum (EUAA). (2022). Study on Language Assessment for Determination of Origin of Applicants for International Protection. https://euaa.europa.eu/sites/default/files/publications/2022-09/Study_on_Language_Assessment_for_Determination_of_Origin_Executive_Summary.pdf

European Union Agency for Fundamental Rights (FRA). (2022). *Bias in Algorithms. Artificial Intelligence and Discrimination*. <https://fra.europa.eu/en/publication/2022/bias-algorithm>

European Union Agency for Fundamental Rights (FRA). (2019). Data quality and artificial intelligence – mitigating bias and error to protect fundamental rights. https://fra.europa.eu/sites/default/files/fra_uploads/fra-2019-data-quality-and-ai_en.pdf

Europol (2025). AI bias in law enforcement - A practical guide, Europol Innovation Lab observatory report. <https://www.europol.europa.eu/publications-events/publications/ai-bias-in-law-enforcement>

Fassier, E. (2021). Prison Phone Companies Are Recording Attorney-Client Calls Across the US. <https://www.vice.com/en/article/prison-phone-companies-are-recording-attorney-client-calls-across-the-us/>

Gallagher R., Jona L. (26 July, 2019). We tested Europe’s new lie detector for travelers – and immediately triggered a false positive. *The Intercept*. <https://theintercept.com/2019/07/26/europe-border-control-ai-lie-detector/>

Garcia Alvarez, G., Tol, R. (2023). The impact of the Bono Social de Electricidad on energy poverty in Spain. University of Sussex. <https://hdl.handle.net/10779/uos.23483840.v1>

García, C., Maqueda, A., Cabo D., Laursen L. (2025). Spanish National Police stop using Veripol, its star AI for detecting false reports. <https://civio.es/transparencia/2025/03/25/national-police-stop-using-veripol-its-star-ai-for-detecting-false-reports/>

Golobardes, M. A., Grasso G., Iamiceli, P. Et. al. (2024). JuLIA Handbook. Artificial Intelligence, Juridicial Decision-Making and Fundamental Rights. https://ssm-italia.eu/wp-content/uploads/2025/02/JuLIA_handbook-Justice_final.pdf

Artificial Intelligence, Judicial Decision-Making and Fundamental Rights. https://ssm-italia.eu/wp-content/uploads/2025/02/JuLIA_handbook-Justice_final.pdf

Kumari, A., Gupta, R., Tanwar, S. et.al. (2020). Blockchain and AI amalgamation for energy cloud management: Challenges, solutions, and future directions. *J. Parallel Distrib. Comput.* 143, p. 148–166. <https://doi.org/10.1016/j.jpdc.2020.05.004>

LaChapelle C., Tucker C. (2023). Generative AI in Political Advertising. The Brennan Center for Justice. <https://www.brennancenter.org/our-work/research-reports/generative-ai-political-advertising>

Latvijas Republikas Tiesībsargs. (2025). Pētījums par diskriminācijas situāciju Latvijā. https://www.tiesibsargs.lv/wp-content/uploads/2025/04/diskriminācijas_situacija_latvija.pdf

Martínez Garay, L., Boix Palop, A., Briz Redón, Á. et.al. (2022). Three predictive policing approaches in Spain: Viogén, RisCanvi and Veripol. Assessment from a human rights perspective. https://www.uv.es/regulation/papers/MartinezGaray2024_Predictive_policing_Spain.pdf

Mazzarella J. (2024). AI in Critical Infrastructure Markets: Are Smart Systems AI? The EU Act Says They May Be. AI Alert. <https://www.mccarter.com/insights/ai-in-critical-infrastructure-markets-are-smart-systems-ai-the-eu-ai-act-says-it-maybe/>

Mccray R. (2025). From Surveillance to Robot Guards: How AI Could Reshape Prison Life. <https://www.themarshallproject.org/2025/08/30/ai-california-prison-machine-surveillance>

Mengidis, N., Tsirikas, T., Vrochidis, S. et.al. (2019). Blockchain and AI for the next generation energy grids: Cyber-security challenges and opportunities. *Information & Security*, 43 (1), p. 21–33. https://isij.eu/system/files/download-count/2023-01/4302-blockchain_ai_energy_grids.pdf

Murad M. (2022). Fighting Back on Algorithmic Opacity. <https://towardsdatascience.com/fighting-back-on-algorithmic-opacity-30a0c13f0224/>

Nguyen, V.N.; Tarekko, W.; Sharma, P. et.al. (2024). Potential of explainable artificial intelligence in advancing renewable energy: Challenges and prospects. *Energy Fuels*, 38, p. 1692–1712. <https://doi.org/10.1021/acs.energyfuels.3c04343>

Orakhelashvili, A. (2024). Real World Cases of Annex III AI Act applications that posed risks to fundamental rights. https://blog.bham.ac.uk/lawresearch/2024/10/real-world-cases-of-annex-iii-ai-act-applications-that-posed-risks-to-fundamental-rights/#_ftn11

Ozkul, D. (2023). Automating Immigration and Asylum: The Uses of New Technologies in Migration and Asylum Governance in Europe. <https://www.rsc.ox.ac.uk/publications/automating-immigration-and-asylum-the-uses-of-new-technologies-in-migration-and-asylum-governance-in-europe>

Park, C. (2025). Addressing Challenges for the Effective Adoption of Artificial Intelligence in the Energy Sector. *Sustainability*, 17(13), 5764. <https://doi.org/10.3390/su17135764>

Pina M.C. (2025). AI in the context of border management, migration and asylum in the EU: technological innovation vs. fundamental rights of migrants in the AI Act. https://officialblogofunio.com/2025/02/15/ai-in-the-context-of-border-management-migration-and-asylum-in-the-eu-technological-innovation-vs-fundamental-rights-of-migrants-in-the-ai-act/?utm_source=chatgpt.com

Shi, Z.; Yao, W.; Li, Z.; Zeng, L. et.al. (2020). Artificial intelligence techniques for stability analysis and control in smart grids: Methodologies, applications, challenges and future directions. *Appl. Energy*, 278. <https://ideas.repec.org/a/eee/appene/v278y2020ics0306261920312228.html>

Sobrinho-García, I. (2021). Artificial Intelligence Risks and Challenges in the Spanish Public Administration: An Exploratory Analysis through Expert Judgements. *Administrative Sciences*. 11(3). <https://doi.org/10.3390/admsci11030102>

Späth, E. (2025). AI Use in the Asylum Procedure in Germany: Exploring Perspectives with Refugees and Supporters on Assessment Criteria and Beyond. In: Ahrweiler, P. (eds)

Participatory Artificial Intelligence in Public Social Services. Artificial Intelligence, Simulation and Society. Springer. https://doi.org/10.1007/978-3-031-71678-2_6

Stankovich, M, Feldfeber, I.; Quiroga Y. (2023). Global Toolkit on AI and the Rule of Law for the Judiciary. <https://www.gcedclearinghouse.org/resources/global-toolkit-ai-and-rule-law-judiciary>

Stolton, S. (8 February, 2021). Commission under pressure in EU court over 'lie detector tech'. *Euractiv*. <https://www.euractiv.com/section/digital/news/aommission-under-pressure-over-lie-detector-tech-in-eu-courts/>

Strielkowski, W, Vlasov, A, Selivanov K, et.al. Prospects and Challenges of the Machine Learning and Data-Driven Methods for the Predictive Analysis of Power Systems: A Review. *Energies*. 2023; 16(10), 4025. <https://doi.org/10.3390/en16104025>

UNESCO. (2023). Global toolkit on AI and the rule of law for the judiciary. <https://unesdoc.unesco.org/ark:/48223/pf0000387331>

US Department of Homeland Security. (2024). Mitigating Artificial Intelligence (AI) Risk: Safety and Security Guidelines for Critical Infrastructure Owners and Operators. https://www.dhs.gov/sites/default/files/2024-04/24_0426_dhs_ai-ci-safety-security-guidelines-508c.pdf

Yeung, K., Harken, A. (2023). How do 'technical' design-choices made when building algorithmic decision-making tools for criminal justice authorities create constitutional dangers? Part II, *Public Law*, <https://arxiv.org/pdf/2301.04715>

Judikatūra

Eiropas Savienības Tiesas 2022. gada 21. jūnija spriedums lietā C-817/19 *Ligue des droits humains*, CLI:EU:C:2022:491

Vispārējās Tiesas 2021. gada 15. decembra spriedums lietā T-158/19 Patrick Breyer pret Eiropas Pētniecības izpildaģentūru, ECLI:EU:T:2021:902