

**მონაცემთა მინიმიზაციისა და მოდელის მაქსიმიზაციის თავსებადობა:
მარეგულირებელი და ეთიკური წინააღმდეგობები ხელოვნური
ინტელექტის განვითარებაში****

ფართომასშტაბიანი ხელოვნური ინტელექტის (AI) სისტემების, განსაკუთრებით კი დიდი ენობრივი მოდელების (LLM) სწრაფმა განვითარებამ, მონაცემთა დაცვის სფეროში დრმა მარეგულირებელი წინააღმდეგობები წარმოშვა. ამ დისკურსის ცენტრშია კონფლიქტი მონაცემთა მინიმიზაციის პრინციპსა, რომელიც „მონაცემთა დაცვის ძირითად რეგულაციით“ არის განმტკიცებული, და ხელოვნური ინტელექტის მოდელების შემუშავების საფუძვლად არსებულ მონაცემთა ინტენსიურ ლოგიკას შორის. სტატია იკვლევს ამ წინააღმდეგობის სამართლებრივ, პრაქტიკულ და ეთიკურ ასპექტებს მონაცემთა დაცვის საზედამხებელო ორგანოების (DPA) პერსპექტივიდან, ანალიზებს აღსრულების მიმდინარე ტენდენციებს, მარეგულირებელ სახელმძღვანელო მითითებებსა და ევროკავშირის ხელოვნური ინტელექტის აქტის მოსალოდნელ გავლენას. სტატიაში წარმოდგენილია არგუმენტი, რომლის თანახმად, მონაცემთა დაცვის საზედამხებელო ორგანოები უნდა გასცდნენ ტრადიციული აღსრულების როლს და იქცნენ ეთიკურ ზედამხებელებად და ხელოვნური ინტელექტის მართვის პროაქტიულ კოორდინატორებად ევროპაში, რათა უზრუნველყონ, რომ პრინციპები არ დასუსტდეს ან იგნორირებულ იქნეს ინოვაციის სახელით.

საკვანძო სიტყვები: მონაცემთა მინიმიზაცია, ხელოვნური ინტელექტის აქტი, დიდი ენობრივი მოდელები („LLM“), მონაცემთა დაცვის ძირითადი რეგულაცია („GDPR“), ხელოვნური ინტელექტის ეთიკური მმართველობა.

* სამართლის მეცნიერებათა ჰაბილიტირებული დოქტორი, ვარშავის კომინისკის უნივერსიტეტის პროფესორი (დოქტორის ხარისხი, 2000 წ., კრაკოვის იაგელონის უნივერსიტეტი; ჰაბილიტაცია, 2016 წ., პოლონეთის მეცნიერებათა აკადემია). პერსონალურ მონაცემთა დაცვის ოფისის პრეზიდენტის მოადგილე.

** ნაშრომი წარმოადგენს პერსონალურ მონაცემთა დაცვის სამსახურის მასპინძლობით ქ. ბათუმში გამართულ პერსონალურ მონაცემთა დაცვის საზედამხებელო ორგანოთა 33-ე ევროპულ კონფერენციაზე („საგაზაფხულო კონფერენციაზე“) წარდგენილი მოხსენების ტექსტს.

1. შესავალი: მონაცემთა დაცვა ფართომასშტაბიანი ხელოვნური ინტელექტის მოდელების ეპოქაში

ხელოვნური ინტელექტის (“AI”) ტექნოლოგიების მზარდმა ინტეგრაციამ საჯარო სერვისებში, კერძო საწარმოებსა და ყოველდღიურ ცხოვრებაში გაამწვავა დიდი ხნის არსებული დაპირისპირება ინოვაციასა და ფუნდამენტური უფლებების დაცვას შორის. მათ შორის ერთ-ერთი ყველაზე მწვავეა ახლად წარმოშობილი კონფლიქტი “GDPR”-ის ფუძემდებლურ მონაცემთა მინიმიზაციის პრინციპსა და მონაცემთა ინტენსიურ ლოგიკას შორის, რომელიც თანამედროვე “AI” სისტემების, განსაკუთრებით კი დიდი ენობრივი მოდელების (“LLM”) განვითარებას უღევს საფუძვლად.

“GDPR”-ი მონაცემთა მინიმიზაციას მონაცემთა კანონიერი დამუშავების ფუძემდებლურ პრინციპად აცხადებს¹. მე-5 მუხლის პირველი პუნქტის „c“ ქვეპუნქტის თანახმად, პერსონალური მონაცემები უნდა იყოს „ადეკვატური, რელევანტური და შეზღუდული იმით, რაც აუცილებელია“ იმ მიზნებთან მიმართებით, რომელთათვისაც ის მუშავდება. აღნიშნული მანდატი მიზნად ისახავს არასაჭირო პერსონალური მონაცემების შეგროვებისა და გამოყენების შეზღუდვას და ხელს უწყობს ანგარიშვალდებულებას, გამჭვირვალობასა და მონაცემთა სუბიექტების უფლებების პატივისცემას. თუმცა “AI” მოდელების, განსაკუთრებით “LLM”-ების, შემუშავების კონტექსტში ეს პრინციპი მზარდი წნეხის ქვეშ ექცევა.

“LLM”-ები ფუნქციონირებს მასშტაბურობის პრინციპით: რაც უფრო ფართო და მრავალფეროვანია მონაცემთა ბაზა, მით უფრო დახვეწილი და მძლავრი ხდება მოდელი². ამის გამო დვეულოპერები ხშირად მიმართავენ ვებსკრაპინგის განურჩევლად გამოყენებას, რათა შეაგროვონ უზარმაზარი მონაცემები, მათ შორის პერსონალური მონაცემები საჯარო ინტერნეტის მეშვეობით. წამყვანი ვარაუდი არის ის, რომ მეტი მონაცემი ნიშნავს მოდელის უკეთეს სიზუსტეს, კონტექსტურ გააზრებასა და ადაპტირების უნარს. თუმცა ეს ვარაუდი ქმნის სტრუქტურულ წინააღმდეგობას მონაცემთა მინიმიზაციასთან, რადგან ასეთი მოდელები შექმნილია არა შეტანილი მონაცემების მინიმიზაციისთვის, არამედ ინფორმაციის მოპოვებისა და განზოგადების შესაძლებლობების მაქსიმიზაციისთვის.

აღსანიშნავია, რომ წამყვანი LLM-ების უმეტესობა შემუშავებულია იმ ორგანიზაციების მიერ, რომელთა სათავე ოფისები ევროკავშირის ფარგლებს გარეთ მდებარეობს. ეს ქმნის დამატებით გამოწვევებს იურისდიქციას, აღსრულებასა და GDPR-ის ექსტრატერიტორიულ გამოყენებასთან დაკავშირებით. ევროპელი მომხმარებლების პერსონალურ მონაცემებს შესაძლოა ამუშავებდნენ არაევროკავშირის სუბიექტები, რომლებიც სრულად

¹ Kuner C., Bygrave L. A., Docksey C., The EU General Data Protection Regulation (GDPR): A commentary. Oxford University Press, 2020.

² Vaezi A., Legal Challenges in the Deployment of Large Language Models: A Comparative Analysis under the GDPR and EU AI Act, 2025.

არ ითავისებენ ევროკავშირის კანონმდებლობით დადგენილ ნორმატიულ და სამართლებრივ ვალდებულებებს. შედეგად, ევროპის მონაცემთა დაცვის საზედამხებელო ორგანოები (“DPA”) დგანან მზარდი აუცილებლობის წინაშე, დაამკვიდრონ ევროკავშირის მონაცემთა დაცვის პრინციპების შესაბამისობა გლობალურ ტექნოლოგიურ კონტექსტში.

ამ გამოწვევების საპასუხოდ, მონაცემთა დაცვის ევროპულმა საბჭომ (“EDPB”) სულ უფრო მეტად იხრება DPA-ების პასუხისმგებლობების უფრო ფართო ინტერპრეტაციისკენ. 2024 წლის ივლისში გამოქვეყნებულ თავის მე-3/2024 განცხადებაში “EDPB”-მ განმარტა, რომ DPA-ები არ არიან მხოლოდ რეაგირებაზე ორიენტირებული მარეგულირებლები, არამედ ასევე პროაქტიული მრჩეველები, კოორდინატორები და ეთიკური არბიტრები მომავალი ხელოვნური ინტელექტის აქტის ფარგლებში³. პასუხისმგებლობების ასეთი გადააზრება ხაზს უსვამს DPA-ების საქმიანობას, ჩაერთონ არა მხოლოდ აღსრულებაში, არამედ “AI” ტექნოლოგიების სტრატეგიულ მმართველობასა და პრევენციულ ზედამხედველობაში.

წინამდებარე სტატია მიზნად ისახავს, შეისწავლოს ამ ცვალებადი მარეგულირებელი ლანდშაფტის შედეგები, განსაკუთრებული აქცენტით მონაცემთა მინიმიზაციასა და მოდელის მაქსიმიზაციას შორის არსებულ წინააღმდეგობებზე. სტატია აანალიზებს, თუ როგორ შეუძლიათ DPA-ებს ეფექტიანად დაიცვან მონაცემთა დაცვის პრინციპები ეპოქაში, რომელშიც მონაცემთა მოცულობა და არა მონაცემთა დისციპლინა, სულ უფრო ხშირად განიხილება ტექნოლოგიური წარმატების საზომად. სამართლებრივი ანალიზის, მარეგულირებელი ნორმების ინტერპრეტაციისა და აღსრულების ახალი პრაქტიკის გათვალისწინების გზით ეს კვლევა ხელს უწყობს პასუხისმგებლიანი ხელოვნური ინტელექტის მომავლის შესახებ აქტუალურ დისკუსიას ევროკავშირისა და მის ფარგლებს გარეთ.

2. სტრუქტურული შეუთავსებლობა მონაცემთა მინიმიზაციასა და ხელოვნური ინტელექტის მოდელის შემუშავებას შორის

ფართომასშტაბიანი ხელოვნური ინტელექტის სისტემების სწრაფმა ევოლუციამ წინა პლანზე წამოწია თანამედროვე მონაცემთა დაცვის სამართლის ფუნდამენტური კითხვა: ვართ თუ არა სამართლებრივ პრინციპებსა და ტექნოლოგიურ პრაქტიკას შორის გარდაუვალი კონფლიქტის მომსწრენი? ამ დილემის საფუძველში დევს ფუნდამენტური დაპირისპირება “GDPR”-ის მონაცემთა მინიმიზაციის პრინციპსა და თანამედროვე მანქანური სწავლების მოდელის, განსაკუთრებით კი დიდი ენობრივი მოდელის (“LLM”) ფუნქციონირების არქიტექტურას შორის.

³ European Data Protection Board, Statement 3/2024 on the role of DPAs under the AI Act, 2024, <https://www.edpb.europa.eu/our-work-tools/our-documents/statements/statement-32024-data-protection-authorities-role-artificial_en> [23.07.2025].

“GDPR”-ი, კერძოდ, მისი მე-5 მუხლის პირველი პუნქტის „c“ ქვეპუნქტი, მონაცემთა მინიმიზაციის პრინციპს მონაცემთა კანონიერი დამუშავების ქვაკუთხედად აცხადებს⁴. ეს პრინციპი მოითხოვს, რომ შეგროვებული პერსონალური მონაცემები იყოს ადეკვატური, რელევანტური და მკაცრად აუცილებელი მოცულობით შემოიფარგლებოდეს იმ კონკრეტული მიზნებისათვის, რომელთათვისაც ის მუშავდება. პრაქტიკაში ეს ნიშნავს, რომ ორგანიზაციებმა თავი უნდა შეიკავონ მონაცემთა შეგროვებისგან ან შენახვისგან, თუკი ეს არ არის გამართლებული მიზნობრიობისა და აუცილებლობის თვალსაზრისით.

ამისგან მკვეთრად განსხვავებით, “LLM”-ის შემუშავების საფუძვლად გამოყენებული ლოგიკა მონაცემთა სიმრავლეს ემყარება. ხელოვნური ინტელექტის დეველოპერებს შორის გაბატონებული მოსაზრების თანახმად, ამ მოდელების ეფექტიანობა და განზოგადების უნარი უმჯობესდება სასწავლო მონაცემების მოცულობასა და მრავალფეროვნებასთან ერთად. შედეგად, “LLM”-ები, როგორც წესი, სწავლობენ უზარმაზარ მონაცემთა ბაზებზე დაყრდნობით, რომლებიც მილიარდობით ტექსტურ ნიმუშს მოიცავს — აკადემიური პუბლიკაციებიდან დაწყებული, ფორუმების პოსტებით, ბლოგებით, სოციალური მედიის კონტენტითა და სხვა საჯაროდ ხელმისაწვდომი წყაროებით დამთავრებული. სასურველი მიზანია ყოვლისმომცველი ლინგვისტური დაფარვისა და სემანტიკური სიმდიდრის მიღწევა, თუმცა ეს მონაცემებზე ორიენტირებული ფილოსოფია პირდაპირ ეწინააღმდეგება “GDPR”-ით დაწესებულ აუცილებლობისა და პროპორციულობის შეზღუდვებს.

ეს სტრუქტურული კონფლიქტი რამდენიმე კრიტიკული მიმართულებით ვლინდება. პირველ რიგში, ვებ-სკრაპინგის განუჩეხლად გამოყენების პრაქტიკას, ხშირად აკლია მკაფიოდ განსაზღვრული მიზანი, რომელიც მონაცემთა მინიმიზაციასთან იქნებოდა თავსებადი. მხოლოდ ვარაუდი, რომ ყველა ხელმისაწვდომმა ტექსტურმა მონაცემმა შეიძლება შეუწყოს ხელი მოდელის გაუმჯობესებას, არასაკმარისია ევროკავშირის მონაცემთა დაცვის კანონმდებლობის მიხედვით, რომელიც კონკრეტულ და ლეგიტიმური დამუშავების მიზნებს მოითხოვს. მეტიც, შეგროვების მასშტაბი, როგორც წესი, მნიშვნელოვნად აღემატება იმას, რაც აუცილებლად ჩაითვლებოდა მოდელის დეკლარირებული ამოცანებისთვის, განსაკუთრებით მაშინ, როდესაც საქმე პერსონალურ მონაცემებს ეხება.

მონაცემთა დაცვის ევროპულმა საბჭომ (EDPB) თავის 28/2024 დასკვნაში ხაზი გაუსვა აუცილებლობის მკაცრი შეფასების მნიშვნელობას. საბჭო მიიჩნევს, რომ სასწავლო მონაცემთა ბაზებში ჩადებული პერსონალური მონაცემები მაშინაც კი, როდესაც მათი პირდაპირი იდენტიფიცირება შეუძლებელია, შეიძლება გახდეს ხელახალი იდენტიფიკაციის რისკების

⁴ Kuner C., Bygrave L. A., Docksey C., *The EU General Data Protection Regulation (GDPR): A commentary*. Oxford University Press, 2020.

საფუძველი⁵. ღრმა სწავლების მოდელებს, თავიანთი არქიტექტურიდან გამომდინარე, შეუძლიათ დაიმახსოვრონ და სიტყვასიტყვით ან პერიფრაზირებული ფორმით აღადგინონ სასწავლო მონაცემები. ეს ქმნის ფარულ საფრთხეს, რომ პერსონალური ინფორმაცია, რომელიც ერთხელ მოხვდა სასწავლო კორპუსში, შესაძლოა მოგვიანებით გამჟღავნდეს მოდელის შედეგების მეშვეობით, დეველოპერის მხრიდან შეგნებული განზრახვის არარსებობის შემთხვევაშიც კი.

“EDPB“-მ ასევე გაამახვილა ყურადღება ანონიმიზაციის შესახებ ზოგადი მტკიცებების არასაკმარისობაზე. მტკიცებულებები იმის შესახებ, რომ პერსონალური მონაცემები საკმარისად დეიდენტიფიცირებული ან ფსევდონიმიზებულია, დასაბუთებული უნდა იყოს კონკრეტული მოდელის შეფასებით და არ შეიძლება ეფუძნებოდეს აბსტრაქტულ ტექნიკურ ვარაუდებს. ინფერენს შეტევებმა, მოდელის ინვერსიამ და წევრობის ინფერენსის ტექნიკებმა ცხადყო, რომ გარკვეულ პირობებში ანონიმიზებული მონაცემები შესაძლოა რევერს-ინჟინერინგს დაექვემდებაროს ან მოხდეს მათი დაკავშირება კონკრეტულ პირებთან. ეს რისკები მოითხოვს მოდელის სწავლებისას გამოყენებული ტექნიკური უსაფრთხოების მექანიზმების ფრთხილ, ინდივიდუალურ ანალიზს.

გარდა ამისა, პროპორციულობის პრინციპი მონაცემთა დამუშავების კანონიერების ცენტრალური ელემენტია. დამუშავებისთვის პასუხისმგებელმა პირებმა უნდა დაასაბუთონ კავშირი დამუშავებული პერსონალური მონაცემების რაოდენობასა და “AI” სისტემის ეფექტიანობით მისაღებ სარგებელს შორის. თუმცა “LLM”-ების შემთხვევაში აუცილებლობისა და პროპორციულობის საზღვრები ხშირად ბუნდოვანია. დეველოპერები ხშირად ვერ ასაბუთებენ, თუ რატომ არის საჭირო მონაცემთა ბაზის კონკრეტული ზომა ან შემადგენლობა, ან რატომ არ განიხილეს უფრო მიზნობრივი და კონფიდენციალურობის დამცავი ალტერნატივები.

შესაბამისად, კონფლიქტი არა უბრალოდ თეორიული, არამედ ღრმად პრაქტიკულია. “AI” დეველოპერები მოქმედებენ პარადიგმაში, რომელიც მონაცემთა მაქსიმალურ ათვისებას ახალისებს მაშინ, როცა მონაცემთა დაცვის ჩარჩოები მოითხოვს შეზღუდვას, დასაბუთებასა და მომხმარებელზე ორიენტირებულ უსაფრთხოების მექანიზმებს. ამ ხარვეზის ამოსავსებად საჭირო იქნება არა მხოლოდ რეგულატორული სიცხადე და აღსრულება, არამედ პარადიგმის ცვლილებაც ხელოვნური ინტელექტის სფეროში ინოვაციების კონცეპტუალიზაციის პროცესში.

თავსებადობის მისაღწევად, “AI”-ის შემუშავების პროცესში სულ უფრო მეტად უნდა მოხდეს „შექმნის ეტაპზე კონფიდენციალურობის უზრუნველყოფისა“ და „ნაგულისხმევი პარამეტრებით კონფიდენციალურობის დაცვის“ პრინციპების ინტეგრირება, რომლებიც GDPR-ის 25-ე მუხლით არის გათვალისწინებული. ეს გულისხმობს მონაცემთა

⁵ European Data Protection Board, Opinion 28/2024 on Training Data for LLMs, 2024 <https://www.edpb.europa.eu/our-work-tools/our-documents/opinion-board-art-64/opinion-282024-certain-data-protection-aspects_en> [23.07.2025].

მინიმიზაციის ლოგიკის ჩაშენებას “AI” სისტემების არქიტექტურაში მათი შექმნის საწყისი ეტაპიდანვე. ეს ასევე ნიშნავს ისეთი მეთოდოლოგიების დანერგვას, რომლებიც ამცირებს პერსონალურ მონაცემებზე დამოკიდებულებას — როგორებიცაა სინთეტიკური მონაცემების გენერირება, ფედერაციული სწავლება ან დიფერენციალური კონფიდენციალურობის მექანიზმები — რითაც ტექნოლოგიური პროგრესი სამართლებრივ ვალდებულებებს შეესაბამება.

დასასრულს, მონაცემთა მინიმიზაციასა და მოდელის მაქსიმიზაციას შორის აღქმული დიქტომია ციფრულ ეპოქაში არსებული უფრო ფართო მმართველობითი გამოწვევების სიმბოლოა. მიუხედავად იმისა, რომ თავისი არსით შეუთავსებელი არ არის, ეს ურთიერთსაწინააღმდეგო ლოგიკები მიზანმიმართულ შეთანხმებას საჭიროებს მულტიდისციპლინური თანამშრომლობის, ტექნიკური ინოვაციებისა და მარეგულირებელი სიფხიზლის გზით. ამ ძალისხმევის გარეშე ფუნდამენტური უფლებების ხელშეუხებლობა შესაძლოა შეილახოს ტექნოლოგიური ოპტიმიზაციისკენ უკონტროლო სწრაფვით.

3. საჯაროდ ხელმისაწვდომი მონაცემები და პერსონალური ინფორმაციის სამართლებრივი სტატუსი “GDPR”-ის მიხედვით

“AI” დეველოპერებს შორის ფართოდ გავრცელებული მოსაზრების თანახმად, საჯაროდ ხელმისაწვდომი მონაცემები თავისი არსით თავისუფლდება მონაცემთა დაცვის რეგულაციების სრული მოქმედებისგან. ეს მცდარი წარმოდგენა საფუძვლად უდევს ბევრ არგუმენტს, რომლებიც ამართლებს ფართომასშტაბიანი მონაცემთა ბაზების შეგროვებასა და გამოყენებას “AI” მოდელების, მათ შორის, დიდი ენობრივი მოდელების (“LLM”) შესასწავლად. დეველოპერები ხშირად ამტკიცებენ, რომ ვინაიდან გამოყენებული მონაცემები უკვე ხელმისაწვდომია ღია ინტერნეტში, მასზე მონაცემთა დაცვის ძირითადი რეგულაციის (“GDPR”)⁶ მარეგულირებელი დაცვის მექანიზმები არ ვრცელდება. თუმცა ასეთი მსჯელობა სამართლებრივად და ეთიკურად ხარვეზიანია.

“GDPR”-ის მიხედვით, მონაცემის „პერსონალურ“ სტატუსს განსაზღვრავს მისი იდენტიფიცირებადობა და არა მისი ხელმისაწვდომობა. “GDPR”-ის მე-4 მუხლის პირველი პუნქტი განმარტავს პერსონალურ მონაცემებს, როგორც ნებისმიერ ინფორმაციას, რომელიც ეხება იდენტიფიცირებულ ან იდენტიფიცირებად ფიზიკურ პირს, მიუხედავად იმისა, ეს ინფორმაცია კერძო თუ საჯარო წყაროებიდან იქნა მოპოვებული. შესაბამისად, ის ფაქტი, რომ

⁶ პოლონეთის მონაცემთა დაცვის საზედამხებელო ორგანოსთან შეხვედრის დროს წარმოდგენილი მოსაზრებები. “OpenAI”-თან ან “Microsoft”-თან შეხვედრების შესახებ დამატებითი ინფორმაციისთვის იხილეთ <www.uodo.gov.pl>.

პერსონალური მონაცემები განთავსებულია საჯარო ფორუმებზე, სოციალურ მედიაში, კომენტარების სექციებში ან სხვა ღია წვდომის დომენებში, არ ართმევს მას დაცულ სტატუსს.

ამ სამართლებრივ პოზიციას მნიშვნელოვანი გავლენა აქვს “AI” სისტემების შემუშავებაზე, განსაკუთრებით იმ სისტემების, რომლებიც ინტერნეტიდან განურჩევლად შეგროვებულ მონაცემებზე სწავლობენ. მონაცემთა დაცვის ევროპულმა საბჭომ (“EDPB”) ამ საკითხთან დაკავშირებით შემოთავაზა გამოხატა რამდენიმე უახლეს დასკვნაში, განსაკუთრებით კი 28/2024 დასკვნაში. “EDPB” ხაზს უსვამს, რომ ასეთი მონაცემების “AI” მოდელების სასწავლებლად გამოყენების კანონიერება უნდა შეფასდეს ყოველი კონკრეტული შემთხვევის მიხედვით. ერთ-ერთი მთავარი საკითხია ის, რომ მონაცემთა ანონიმიზაციის პროცესი — რომელსაც დევლოპერები ხშირად შესაბამისობის უზრუნველყოფის ღონისძიებად ასახელებენ — არ შეიძლება ეფექტიანად ჩაითვალოს მკაცრი, კონკრეტული მოდელისთვის დამახასიათებელი ვალიდაციის გარეშე.

მაშინაც კი, თუ “AI” მოდელი არ არის შექმნილი პერსონალური მონაცემების პირდაპირ გამოსაცემად, რჩება მნიშვნელოვანი რისკი იმისა, რომ სასწავლო კორპუსში ჩადებული პერსონალური ინფორმაცია შენარჩუნებული იყოს მოდელის პარამეტრებში. ეს ფარული მონაცემები შესაძლოა უნებლიეთ აღდგეს მომხმარებლის მოთხოვნების საპასუხოდ, რითაც შეიქმნება კონფიდენციალურობის დარღვევის პოტენციალი დასკვნის გამოტანის ან ხელახალი იდენტიფიკაციის გზით. უახლესმა აკადემიურმა კვლევებმა და რეალურმა ინციდენტებმა აჩვენა, რომ “LLM”-ებს შეუძლიათ უნებლიეთ აღადგინონ კონკრეტული პერსონალური დეტალები, როგორებიცაა სახელები, მისამართები ან პირადი საუბრების ფრაგმენტები, რაც სერიოზულ კითხვებს აჩენს “AI” კონტექსტში სტანდარტული ანონიმიზაციის ტექნიკის საკმარისობასთან დაკავშირებით.

ამ რისკების ფონზე “DPA”-ები სულ უფრო მეტად ამოწმებენ ანონიმიზაციის შესახებ მტკიცებულებებს და ითხოვენ გამჭვირვალობას სასწავლო მონაცემთა შემადგენლობასთან დაკავშირებით. შესაბამისობის სათანადოდ შესაფასებლად მარეგულირებლებს უნდა ჰქონდეთ წვდომა დეტალურ ტექნიკურ დოკუმენტაციაზე, მათ შორის, მონაცემთა ბაზების წყაროებზე, წინასწარი დამუშავების მეთოდებსა და რისკების შემარბილებელ სტრატეგიებზე. ეს აძლიერებს მარეგულირებელი ორგანოების საჭიროებას, ინვესტიცია ჩადონ ტექნიკურ ექსპერტიზასა და მრავალდისციპლინური პოტენციალის განვითარებაში. ასეთი შესაძლებლობების გარეშე “DPA”-ები ვერ შეძლებენ იმ დეტალური შეფასებების ჩატარებას, რომლებიც აუცილებელია იმის დასადგენად, დაკმაყოფილებულია თუ არა მონაცემთა მინიმიზაციის, აუცილებლობისა და პროპორციულობის სტანდარტები.

გარდა ამისა, „შექმნის ეტაპზე კონფიდენციალურობის უზრუნველყოფის პრინციპი“, რომელიც “GDPR”-ის 25-ე მუხლშია ჩამოყალიბებული, მოითხოვს, რომ მონაცემთა დაცვის გარანტიები თავიდანვე იყოს ჩაშენებული დამუშავების პროცესებში. ეს გულისხმობს მონაცემთა დაცვაზე ზეგავლენის

შეფასების ("DPIA") საფუძვლიანად ჩატარებას სასწავლო ოპერაციების დაწყებამდე, მონაცემთა შეგროვების მიზნის, მასშტაბისა და შეზღუდვების განსაზღვრით. დეველოპერებმა უნდა განსაზღვრონ პერსონალური მონაცემების რომელი კატეგორიებია აუცილებელი მოდელის მიზნების მისაღწევად და დაასაბუთონ, რომ განიხილეს ნაკლებად ინვაზიური ალტერნატივები.

თანაბრად მნიშვნელოვანია პროპორციულობის კონცეფცია, რომელიც მოითხოვს დასაბუთებულ კავშირს დამუშავებული პერსონალური მონაცემების რაოდენობას, სენსიტიურობასა და დასახულ ლეგიტიმურ მიზნებს შორის. ონლაინ კონტენტის მასობრივი და განურჩეველი სკრაპინგი, განსაკუთრებით კონტექსტუალური ფილტრაციის ან თანხმობის გარეშე, არსებითად ეჭვქვეშ აყენებს პროპორციულობის პრინციპს და ძირს უთხრის მომხმარებლის ნდობას ციფრული ეკოსისტემებისადმი.

იდეა, რომ „საჯარო ნიშნავს დასაშვებს“, ცალსახად უნდა იქნეს უარყოფილი. ის ფაქტი, რომ ინფორმაცია ხელმისაწვდომია ონლაინ, არ ანიჭებს ლიცენზიას მისი მიზნის შესაცვლელად და მანქანური სწავლებისთვის გამოსაყენებლად შესაბამისი სამართლებრივი და ეთიკური გარანტიების გარეშე. "DPA"-ების მოვალეობაა, ეჭვქვეშ დააყენონ ეს ნორმა და გაამყარონ განსხვავება ხილვადობასა და ვალიდურობას შორის მონაცემთა მართვაში.

დასასრულს, საჯარო მონაცემების კანონიერი გამოყენება "AI"-ის შემუშავების პროცესში ჯერ კიდევ გადაუჭრელი საკითხია. ის მოითხოვს მყარ სამართლებრივ ინტერპრეტაციას, მკაცრ ტექნიკურ შემოწმებასა და პროაქტიულ მარეგულირებელ პოზიციას, რათა უზრუნველყოფილ იქნას, რომ ინდივიდუალური უფლებები არ დაექვემდებაროს ტექნოლოგიური ექსპანსიის იმპერატივებს.

4. დამუშავების სამართლებრივი საფუძვლები "AI" მოდელების სწავლებისას: თანხმობა, ლეგიტიმური ინტერესები და გამჭვირვალობის სირთულე

"AI" მოდელების, განსაკუთრებით კი დიდი ენობრივი მოდელების ("LLM"), სწავლებისას გამოყენებული პერსონალური მონაცემების დამუშავებისთვის ვალიდური სამართლებრივი საფუძვლის დადგენა ერთ-ერთი ყველაზე რთული და სადავო საკითხია არსებულ მარეგულირებელ გარემოში. მიუხედავად იმისა, რომ "AI" შესაძლებლობების გასავითარებლად სულ უფრო მეტად ეყრდნობიან მასიურ მონაცემთა კორპუსებს, ბევრ დეველოპერს არ განუმარტავს, თუ როგორ შეესაბამება ასეთი დამუშავება მონაცემთა დაცვის ძირითად რეგულაციას ("GDPR"), განსაკუთრებით მისი მე-6, მე-7, მე-13 და მე-14 მუხლების გათვალისწინებით.

მიუხედავად იმისა, რომ თანხმობა ხშირად "GDPR"-ით გათვალისწინებული კანონიერი დამუშავების ოქროს სტანდარტად მიიჩნევა,

მისი პრაქტიკული გამოყენება “AI”-ის სწავლების კონტექსტში სირთულეებითაა სავსე. მე-4 მუხლის მე-11 პუნქტი თანხმობას განმარტავს, როგორც მონაცემთა სუბიექტის სურვილების თავისუფლად გამოხატულ, კონკრეტულ, ინფორმირებულ და ცალსახა მითითებას. თუმცა, ეს სტანდარტი იშვიათად სრულდება, როდესაც თანხმობის მიღება ხდება ზოგადი მომსახურების პირობებით, გაუმჭვირვალე კონფიდენციალურობის პოლიტიკით ან პლატფორმის დონის შეტყობინებებით. როდესაც პერსონალური მონაცემები შეგროვებულია საჯარო ვებსაიტებიდან ან მომხმარებლის მიერ გენერირებული კონტენტის პლატფორმებიდან, მონაცემთა სუბიექტებს, როგორც წესი, არ აქვთ რეალური შესაძლებლობა, მისცენ თანხმობა ან თქვან მასზე უარი, რომ აღარაფერი ითქვას იმის გაგებაზე, თუ როგორ იქნება მათი მონაცემები ხელახლა გამოყენებული “AI” სისტემებში. შესაბამისად, ამ კონტექსტში თანხმობაზე დაყრდნობა ხშირად იურიდიულად არასაკმარისი და ეთიკურად საეჭვოა.

პრაქტიკაში ბევრი დევლოპერი, როგორც უფრო მოქნილ ალტერნატივას, მე-6 მუხლის პირველი პუნქტის „f“ ქვეპუნქტით გათვალისწინებულ ლეგიტიმური ინტერესის სამართლებრივ საფუძველს იყენებს⁷. თუმცა ამ საფუძვლის გამოყენების ზღვარი მკაცრია. დამუშავებისთვის პასუხისმგებელმა პირმა უნდა ჩაატაროს ყოვლისმომცველი სამნაწილიანი დაბალანსების ტესტი: (1) განსაზღვროს მონაცემთა დამუშავებისთვის პასუხისმგებელი პირის ან მესამე მხარის მიერ დაცული ლეგიტიმური ინტერესი, (2) დაასაბუთოს, რომ დამუშავება აუცილებელია ამ ინტერესის მისაღწევად, და (3) დაამტკიცოს, რომ ეს ინტერესი არ აღემატება მონაცემთა სუბიექტის უფლებებს და თავისუფლებებს. “EDPB”-ის 2024 წლის სახელმძღვანელო მითითებების თანახმად, რომელიც “AI”-სა და მონაცემთა დამუშავებას ეხება, ეს შეფასება უნდა იყოს ობიექტური და მტკიცებულებებზე დაფუძნებული და გამყარებული უნდა იყოს დოკუმენტირებული უსაფრთხოების გარანტიებით, ანგარიშვალდებულების ზომებითა და მონაცემთა სუბიექტის რისკების შემარბილებელი სტრატეგიებით.

ამ მიდგომის მთავარი სირთულე ის არის, რომ აუცილებლობა და პროპორციულობა იშვიათად არის თავისთავად ცხადი “LLM”-ების შემუშავების კონტექსტში. მონაცემების სკრაპინგის მასშტაბისა და გაუმჭვირვალობის, ასევე, მრავალფეროვანი სასწავლო მონაცემებიდან მიღებული სარგებლის სპეკულაციური ბუნების გათვალისწინებით, ხშირად გაუგებარია, არის თუ არა დამუშავება მკაცრად აუცილებელი დეკლარირებული მიზნისთვის, თუ უბრალოდ მოსახერხებელია მოდელის ეფექტიანობის მაქსიმალურად გასაზრდელად. მტკიცების ტვირთი ეკისრება მონაცემთა დამუშავებისთვის პასუხისმგებელ პირს, რომელმაც უნდა ახსნას, თუ რატომ ვერ მოხერხდა მსგავსი შედეგების მიღწევა ალტერნატიული, ნაკლებად ინვაზიური მეთოდებით.

⁷ Sangaraju V. V., AI and Data Privacy in Healthcare: Compliance with HIPAA, GDPR, and Emerging Regulations, International Journal of Emerging Trends in Computer Science and Information Technology, 2025, 67–74.

ამ გამოწვევებს კიდევ უფრო ართულებს გამჭვირვალობის უზრუნველყოფის ვალდებულება, რომელიც “GDPR”-ის მე-13 და მე-14 მუხლებით არის გათვალისწინებული. ეს დებულებები ავალდებულებს მონაცემთა დამუშავებისთვის პასუხისმგებელ პირებს, აცნობონ მონაცემთა სუბიექტებს მათი პერსონალური მონაცემების შეგროვებისა და გამოყენების შესახებ, მიუხედავად იმისა, ეს მონაცემები მიღებულია პირდაპირ თუ ირიბად. მრავალი პლატფორმიდან ფართომასშტაბიანი სკრაპინგის საფუძველზე “AI”-ის სწავლების კონტექსტში ამ ვალდებულების შესრულება თითქმის შეუძლებელი ხდება. დევლოპერებს იშვიათად აქვთ წვდომა იმ პირთა ვინაობაზე ან საკონტაქტო ინფორმაციაზე, რომელთა მონაცემებიც სასწავლო ბაზებში მოხვდა, ხოლო რეტროაქტიული შეტყობინება ოპერაციულად შეუძლებელია.

მიუხედავად ამისა, “GDPR”-ი არ ითვალისწინებს გამონაკლისს გამჭვირვალობის ვალდებულებებთან დაკავშირებით მასშტაბის ან ტექნიკური არაპრაქტიკულობის საფუძველზე. გამჭვირვალობის ეფექტიანი მექანიზმების არარსებობის პირობებში საფუძვლად მდებარე მონაცემთა დამუშავების კანონიერება ეჭვქვეშ დგება. ამას მნიშვნელოვანი შედეგები მოაქვს იმ დევლოპერებისთვის, რომლებიც ცდილობენ დაეყრდნონ ლეგიტიმურ ინტერესს: თუ არ მოხდება მონაცემთა სუბიექტების ინფორმირება, იზღუდება მათი შესაძლებლობა, განახორციელონ თავიანთი უფლებები, როგორიცაა, მაგალითად, 21-ე მუხლით გათვალისწინებული შეჩინააღმდეგების უფლება, რაც, თავის მხრივ, კიდევ უფრო ასუსტებს მონაცემთა დამუშავების ლეგიტიმურობას.

უფრო მეტიც, გამჭვირვალობა არა მხოლოდ სამართლებრივი მოთხოვნაა, არამედ უმნიშვნელოვანესი ეთიკური და საზოგადოებრივი იმპერატივი. ნდობა “AI” სისტემებისა და მათი მმართველი ინსტიტუტების მიმართ დამოკიდებულია ინდივიდების უნარზე, გაიგონ, როგორ გამოიყენება მათი მონაცემები და შეინარჩუნონ გარკვეული კონტროლი ამ გამოყენებაზე. ბევრი “LLM”-ის გაუმჭვირვალობა როგორც მათი სასწავლო მონაცემების, ისე ფუნქციონირების თვალსაზრისით, კიდევ უფრო აფართოებს ანგარიშვალდებულების ნაკლოვანებას, რაც აფერხებს დემოკრატიულ ზედამხედველობასა და საზოგადოებრივ ნდობას.

დასასრულს, AI-ის სწავლებისთვის ყველაზე ხშირად გამოყენებული სამართლებრივი საფუძვლები — თანხმობა და ლეგიტიმური ინტერესი — პრაქტიკაში უადრესად პრობლემურია⁸. დევლოპერები უნდა გასცდნენ ზედპირულ შესაბამისობას და გაითავისონ მონაცემთა დაცვის კანონის სულისკვეთება გამჭვირვალობის, ანგარიშვალდებულებისა და აუცილებლობის პრინციპების ჩაშენებით “AI” სისტემების დიზაინსა და დანერგვაში. ამგვარი ძალისხმევის გარეშე აღნიშნულ სამართლებრივ

⁸ Hoofnagle C. J., van der Sloot B., Zuiderveen Borgesius F., The European Union General Data Protection Regulation: What it is and what it means. Information & Communications Technology Law, 28(1), 2019, 65–98.

საფუძვლებზე დაყრდნობამ შესაძლოა, არა მხოლოდ ვერ გაუძლოს მარეგულირებლის შემოწმებას, არამედ შეარყიოს ხელოვნური ინტელექტის განვითარების ლეგიტიმურობა როგორც საზოგადოების, ისე პოლიტიკის შემქმნელების თვალში.

5. აღსრულების ტრაექტორიები და სტრატეგიული მარეგულირებელი რეაგირება “AI”-ის მიერ მონაცემთა გამოყენების პრაქტიკაზე

“GDPR”-ის დებულებების აღსრულება “AI”-ის განვითარების სფეროში თანამედროვე მონაცემთა დაცვის ერთ-ერთი ყველაზე რთული მიმართულებაა. “AI” სისტემების, განსაკუთრებით კი დიდი ენობრივი მოდელების (“LLM”) სირთულე მარეგულირებლებს მრავალმხრივი გამოწვევების წინაშე აყენებს, მათ შორის, ტექნიკური გაუმჭირვალობის, მონაცემთა გლობალიზებული ნაკადებისა და იურიდიული სფეროსი ანგარიშვალდებულების ხარვეზების სახით. ამ ბარიერების მიუხედავად, შეინიშნება შესამჩნევი ევოლუცია ევროპის მონაცემთა დაცვის საზედამხებელო ორგანოების (“DPA”) პოზიციაში, რომლებიც სულ უფრო მეტად გადადიან რეაქციული აღსრულებიდან კოორდინირებულ და პროაქტიულ ზედამხედველობაზე.

აღსრულების ერთ-ერთი უმნიშვნელოვანესი ეტაპი 2023 წელს დადგა, როდესაც იტალიის მონაცემთა დაცვის საზედამხებელო ორგანომ (“Garante per la protezione dei dati personali”) “ChatGPT”-ზე დროებითი აკრძალვა დააწესა. გადაწყვეტილება რამდენიმე საფუძველს ეყრდნობოდა, მათ შორის, მონაცემთა დამუშავების კანონიერი საფუძვლის არარსებობას, არასაკმარის გამჭვირვალობასა და იმ მექანიზმების არარსებობას, რომლებიც მონაცემთა სუბიექტებს თავიანთი უფლებებით სარგებლობის საშუალებას მისცემდა. ეს იყო ევროპის მონაცემთა დაცვის საზედამხებელო ორგანოს (“DPA”) მიერ საბაზისო მოდელის წინააღმდეგ განხორციელებული პირველი გახმაურებული ჩარევა, რაც იმის ნიშანია, რომ ხელოვნური ინტელექტის მასშტაბური სისტემებიც კი არ არის დაცული “GDPR”-ის აღსრულებისგან⁹.

სხვა “DPA”-ებიც მიჰყვნენ ამ მაგალითს როგორც აღსრულების, ისე სახელმძღვანელო მითითებების შემუშავების კუთხით. საფრანგეთის მონაცემთა დაცვის ორგანომ (“CNIL”) გამოსცა ყოვლისმომცველი რეკომენდაციები ვებ-სკრაპინგთან დაკავშირებით და ხაზი გაუსვა, რომ საჯარო ხელმისაწვდომობა არ უტოლდება იურიდიულ დასაშვებობას¹⁰. გაერთიანებული სამეფოს ინფორმაციის კომისრის ოფისმა (“ICO”) ანალოგიურად გამოაქვეყნა სახელმძღვანელო მითითებები, რომლებშიც

⁹ Garante per la protezione dei dati personali, Decision on OpenAI (ChatGPT), 2023, <[https://gdprhub.eu/index.php?title=Garante_per_la_protezione_dei_dati_personali_\(Italy\)_-_9870832](https://gdprhub.eu/index.php?title=Garante_per_la_protezione_dei_dati_personali_(Italy)_-_9870832)> [23.07.2025].

¹⁰ CNIL, First Recommendations on the Development of AI Systems, Commission Nationale de l’Informatique et des Libertés, 2024, <<https://www.cnil.fr/en/ai-cnil-publishes-its-first-recommendations-development-artificial-intelligence-systems>> [23.07.2025].

განმარტებულია პირობები, რომელთა დაცვითაც “AI”-ის დეველოპერებს შეუძლიათ კანონიერად გამოიყენონ საჯარო წყაროებიდან მოპოვებული მონაცემები სასწავლო მიზნებისთვის¹¹. ამასობაში საბერძნეთის მონაცემთა დაცვის ორგანომ “Clearview AI” 20 მილიონი ევროთი დააჯარიმა, რითაც “GDPR”-ის დებულებების შესაბამისად, ბიომეტრიული მონაცემების სკრაპინგისთვის პასუხისმგებლობის დაკისრების პრეცედენტი შექმნა¹². ეს ქმედებები ცხადყოფს მონაცემთა დაცვის საზედამხედველო ორგანოების (“DPA”) მზარდ მზაობას, დაუპირისპირდნენ ხელოვნური ინტელექტის სფეროს გავლენიან მოთამაშეებს.

ევროპულ დონეზე მთავარი სიახლე ერთობლივი აღსრულების მექანიზმების გაჩენაა. მონაცემთა დაცვის ევროპულმა საბჭომ (“EDPB”) შექმნა სპეციალური სამუშაო ჯგუფები. განსაკუთრებით აღსანიშნავია “ChatGPT”-ისა და “DeepSeek”-ის¹³ შესახებ შექმნილი ჯგუფები, რომლებიც წევრ სახელმწიფოებში კოორდინირებული გამოძიებებისა და ჰარმონიზებული ინტერპრეტაციების საშუალებას იძლევა. სავარაუდოდ, ეს ძალისხმევა ინსტიტუციონალიზებული იქნება ხელოვნური ინტელექტის აქტის ფარგლებში, რომელიც ევროკავშირის მასშტაბით “AI”-ის მართვის ჩარჩოს ნერგავს და მოიცავს ხელოვნური ინტელექტის ევროპულ ოფისსა და გაძლიერებულ ტრანსსასაზღვრო საზედამხედველო სტრუქტურებს. ეს ახალი მექანიზმები იმეორებს “GDPR”-ის „ერთი ფანჯრის პრინციპის“ მოდელს, თუმცა უფრო ძლიერი აქცენტით სისტემური რისკების შეფასებასა და სექტორულ ზედამხედველობაზე.

პარალელურად, “DPA”-ები იწყებენ სამომავლოდ ორიენტირებული მარეგულირებელი სტრატეგიების შემუშავებას. ეს მოიცავს პროაქტიული სახელმძღვანელო მითითებების გამოცემას, ალგორითმული ზეგავლენის შეფასებების მოთხოვნასა და ტექნიკური აუდიტის პროცედურების შესწავლას მოდელის ახსნადობისა და სასწავლო მონაცემების წარმომავლობის დასადგენად. ტრეექტორია ნათელია: აღსრულება აღარ შემოიფარგლება მხოლოდ წარსულში ჩადენილი დარღვევების დასჯით, არამედ ახლა მოიცავს *ex ante* (წინასწარ) რეგულირებას, რომელიც მიზნად ისახავს სისტემური ზიანის თავიდან აცილებას, სანამ ის რეალობად იქცევა.

6. ეთიკური ზედამხედველობა “AI”-ის განვითარებაში: მონაცემთა დაცვის საზედამხედველო ორგანოების როლის გაფართოება

¹¹ ICO, The Lawful Basis for Web Scraping to Train Generative AI Models, Information Commissioner's Office (UK), 2024, <<https://ico.org.uk/about-the-ico/what-we-do/our-work-on-artificial-intelligence/response-to-the-consultation-series-on-generative-ai/the-lawful-basis-for-web-scraping-to-train-generative-ai-models/>> [23.07.2025].

¹² Info on: <https://www.edpb.europa.eu/news/national-news/2022/hellenic-dpa-fines-clearview-ai-20-million-euros_en> [23.07.2025].

¹³ Deng Z., Ma W., Han Q. L., Zhou W., Zhu X., Exploring DeepSeek: A Survey on Advances, Applications, Challenges and Future Directions, IEEE/CAA Journal of Automatica Sinica, 12(5), 2025, 872–893.

მიუხედავად იმისა, რომ სამართლებრივი ჩარჩოები, როგორიცაა “GDPR” და მომავალი ხელოვნური ინტელექტის აქტი, შესაბამისობის ფორმალურ კრიტერიუმებს აღგენს, მათ ყოველთვის არ შეუძლიათ სრულად მოიცვან “AI”-ის განვითარების ეთიკური განზომილებები. გენერაციული “AI” მოდელების სასწავლებლად პერსონალური მონაცემების გამოყენება წარმოშობს უფრო ფართო საზოგადოებრივ შემფოთებას, რომელიც უკავშირდება ადამიანის ღირსებას, ინდივიდუალურ ავტონომიასა და კულტურულ რეპრეზენტაციას. პრაქტიკა, რომელიც კანონის ვიწრო ინტერპრეტაციით ტექნიკურად კანონიერია, შესაძლოა მაინც იწვევდეს ეთიკურ დისკომფორტს, საზოგადოებრივ უკმაყოფილებას ან სოციალურ ზიანს.

ამის თვალსაჩინო მაგალითია ექსპრესიული პერსონალური მონაცემების, - როგორებიცაა ხმოვანი ჩანაწერები, ბიომეტრიული გამოსახულებები ან შემოქმედებითი კონტენტი, გამოყენება სინთეტიკური მედიის გენერირებისთვის. მიუხედავად იმისა, რომ დეველოპერებმა შესაძლოა განაცხადონ, რომ მონაცემთა ასეთი გამოყენება კანონიერ საფუძვლებს ემყარება, თუკი ეს მონაცემები საჯაროდ ხელმისაწვდომი იყო, ეს მიდგომა უგულებელყოფს ისეთ უფრო ღრმა საკითხებს, როგორებიცაა თანხმობა, შემოქმედებითი საკუთრება და საკუთარი გამოსახულების უფლება. ექსპლუატაციდან გამომდინარე, შემოქმედებით და ჟურნალისტურ სექტორებში საარქივო მასალის უნებართვო გამოყენება “AI”-ის მიერ გენერირებული ხმების ან ავტარების სასწავლებლად ფართოდაა დაგმოხილი. ასეთი პრაქტიკა არღვევს პროფესიულ კეთილსინდისიერებას და ამცირებს ინდივიდების უნარს, გააკონტროლონ, თუ როგორ ხდება მათი იდენტობის ციფრული რეპროდუქცია.

ეს შემფოთების საგანი სცდება ინდივიდუალურ ზიანს და მოიცავს სისტემურ რისკებს, მათ შორის, „დიპფიკებს“, დეზინფორმაციას, პოლიტიკურ მანიპულაციასა და კულტურულ ჰომოგენიზაციას. დიდი ენობრივი მოდელების („LLM“) ან მულტიმოდალური ხელოვნური ინტელექტის სისტემების მიერ გენერირებული სინთეტიკური კონტენტი შესაძლოა გამოყენებულ იქნას რეალური პირების იმიტირების, საზოგადოებრივი დისკურსის დამახინჯების ან დემოკრატიული პროცესებისთვის ძირის გამოთხრის მიზნით. ასეთი გამოყენების ეთიკური შედეგები ღრმაა და მათი წინასწარ განჭვრეტა ყოველთვის არ არის შესაძლებელი მონაცემთა შეგროვების ან მოდელის სწავლების ეტაპზე.

ამ რისკების გათვალისწინებით, “EDPB“-მ მოუწოდა “DPA“-ებს, იკისრონ უფრო ფართო ეთიკური მანდატი. ფორმალური შესაბამისობის უზრუნველყოფის მიღმა, “DPA“-ებისგან სულ უფრო მეტად მოელიან, რომ გამოავლინონ ძალაუფლების ასიმეტრია ინდივიდებსა და დეველოპერებს შორის, შეაფასონ ფართომასშტაბიანი მონაცემთა ექსპლუატაციის საზოგადოებრივი გავლენა და ხელი შეუწყონ პასუხისმგებლიან ინოვაციას. შესაბამისად, ეთიკური შეფასება უნდა განიხილებოდეს, როგორც მონაცემთა დაცვის მმართველობის დამატებითი განზომილება, რომელიც

ინტეგრირებულია რისკზე დაფუძნებულ რეგულირებასა და გამჭვირვალობის ვალდებულებებში.

ამ როლის ეფექტიანად შესასრულებლად DPA-ებმა უნდა განავითარონ ინტერდისციპლინური შესაძლებლობები, რომლებიც მოიცავს იურიდიულ, ტექნიკურ, სოციოლოგიურ და ფილოსოფიურ ექსპერტიზას, და წარმართონ დიალოგი დაზარალებულ თემებთან, სამოქალაქო საზოგადოების ორგანიზაციებთან და აკადემიურ მკვლევარებთან. ეთიკური ზედამხედველობა უნდა იყოს ჩაშენებული ალგორითმულ დიზაინში ისეთი მექანიზმების საშუალებით, როგორებიცაა: სამართლიანობის აუდიტი, მოდელის შეფასებაში მონაწილეობა და საზოგადოებრივი ინტერესების ზეგავლენის შეფასების დოკუმენტები.

შეჯამების სახით, ეთიკური მეურვეობა ყალიბდება, როგორც მარეგულირებელი ფუნქციის კრიტიკული გაფართოება. ის ხელახლა აყალიბებს “DPA”-ების როლს არა მხოლოდ როგორც კანონიერების მცველების, არამედ როგორც სამართლიანობის არბიტრების სწრაფად განვითარებად ტექნოლოგიურ ლანდშაფტში. მხოლოდ სამართლებრივი შესაბამისობის ეთიკურ ლეგიტიმურობასთან შეთავსებით შეძლებს “AI” სისტემების მმართველობა შეინარჩუნოს საზოგადოებრივი ნდობა და ფუნდამენტური უფლებები ციფრულ ეპოქაში.

7. სამომავლო მიმართულებები: ხელოვნური ინტელექტის აქტის ინტეგრაცია და ინსტიტუციური გამოწვევების გადაჭრა

ჩვენ მიერ უკვე ნახსენები ევროკავშირის ხელოვნური ინტელექტის აქტი (“AI Act”) წარმოადგენს უმნიშვნელოვანეს მარეგულირებელ ინიციატივას, რომელიც “AI” სისტემების მართვისთვის ყოვლისმომცველ, რისკზე დაფუძნებულ ჩარჩოს ამკვირდებს. “GDPR”-ის მიერ ჩაყრილ საფუძვლებზე დაყრდნობით, “AI Act” აწესებს მკაცრ ვალდებულებებს ე. წ. მაღალი რისკის “AI” აპლიკაციებისთვის. ეს, სხვა საკითხებთან ერთად, მოიცავს შესაბამისობის სავალდებულო შეფასებებს, სტრუქტურირებულ ტექნიკურ დოკუმენტაციას, მონაცემთა მართვის მყარ ჩარჩოებსა და ბაზარზე განთავსების შემდგომი მონიტორინგის მოთხოვნებს.

“AI Act”-ის ერთ-ერთი ყველაზე მნიშვნელოვანი წვლილი არის მონაცემთა ხარისხისა და მონაცემთა მინიმიზაციის პრინციპების ფუქნციონირება “AI”-ის სასიცოცხლო ციკლის ფარგლებში. “AI Act”-ის მე-10 მუხლის მე-3 პუნქტი ადგენს, რომ სწავლების, ვალიდაციისა და ტესტირებისთვის გამოყენებული მონაცემთა ბაზები უნდა იყოს „რელევანტური, უშეცდომო და მაქსიმალურად სრულყოფილი“. ამასთან, ხაზს უსვამს, რომ მათი მასშტაბი უნდა იყოს მინიმიზებული პერსონალური მონაცემების არასაჭირო გამჟღავნების თავიდან ასაცილებლად. ამ თვალსაზრისით, აქტი აძლიერებს “GDPR”-ის მიერ

დაწყებულ ნორმატიულ ტრაექტორიას მონაცემთა დაცვის სტანდარტების პირდაპირ “AI” სისტემის დიზაინსა და შეფასებაში ჩაშენებით.

მონაცემთა დაცვის საზედამხებდველო ორგანოებისთვის (“DPA”) “AI Act”-ის იმპლემენტაცია ნიშნავს როგორც პასუხისმგებლობების გაფართოებას, ისე ინსტიტუციური იდენტობის ტრანსფორმაციას. “DPA”-ები აღარ იმოქმედებენ მხოლოდ, როგორც მონაცემთა კონფიდენციალურობის ეროვნული აღმსრულებლები, არამედ უნდა ჩამოყალიბდნენ “AI”-ის მართვის პანევროპული ქსელის საკვანძო რგოლებად. ეს მოიცავს თანამშრომლობას ახლად დაარსებულ ხელოვნური ინტელექტის ევროპულ ოფისთან, მონაწილეობას ტრანსსასაზღვრო გამოძიებებში, წვლილის შეტანას ჰარმონიზებული სახელმძღვანელო მითითებების შემუშავებასა და მაღალი რისკის “AI” სისტემების ზედამხედველობაში, რომლებიც გამოიყენება მრავალ სექტორში, მათ შორის, ჯანდაცვაში, განათლებაში, დასაქმებასა და სამართალდაცვაში.

თუმცა აღნიშნული გარდამავალი პერიოდი გამოწვევებითაა სავსე. ბევრი “DPA” ამჟამად რესურსების მნიშვნელოვანი შეზღუდვების წინაშე დგას, მათ შორის, კადრების ნაკლებობის, შეზღუდული ტექნიკური ინფრასტრუქტურისა და მანქანურ სწავლებაში, ალგორითმულ აუდიტსა და სისტემურ ინჟინერიაში არასაკმარისი შიდა ექსპერტიზის სახით. “AI Act”-ით განსაზღვრული ახალი პასუხისმგებლობები, როგორებიცაა: სასწავლო მონაცემების წარმომავლობის შეფასების, ალგორითმული ზეგავლენის შეფასებისა და დიზაინის გამჭვირვალობასთან შესაბამისობის უზრუნველყოფის შესაძლებლობა, მოითხოვს მნიშვნელოვან ინვესტიციას ორგანიზაციულ შესაძლებლობებში, უნარების განვითარებასა და ინსტიტუციურ კოორდინაციაში.

სირთულის კიდევ ერთი წყარო “GDPR”-სა და ხელოვნური ინტელექტის აქტს შორის არსებული დოქტრინული და ოპერაციული კვეთიდან გამომდინარეობს. დეველოპერებმა უნდა იხელმძღვანელონ ერთმანეთის თანმხვედრი და ზოგჯერ პოტენციურად ურთიერთსაწინააღმდეგო ვალდებულებებით, რომლებიც ეხება მონაცემთა მინიმიზაციას, დამუშავების კანონიერ საფუძველს, სამართლიანობას, ანგარიშვალდებულებასა და მონაცემთა სუბიექტის უფლებებს. ეს თანხვედრები მოითხოვს განმარტებით მითითებებს EDPB-სა და ხელოვნური ინტელექტის ევროპული ოფისისგან, რათა უზრუნველყოფილ იქნას თანმიმდევრული აღსრულება. შესაბამისობის ერთობლივი ჩარჩოებისა და სამოდელო შაბლონების შემუშავება ხელს შეუწყობს მარეგულირებელი ხარვეზების აღმოფხვრასა და სამართლებრივი გარკვეულობის გაზრდას იმ დეველოპერებისთვის, რომლებიც ევროკავშირის მრავალ იურისდიქციაში ოპერირებენ.

დაბოლოს, “AI”-ის განვითარების გლობალური ბუნება გამოწვევებს უქმნის ევროპული რეგულაციების აღსრულების მასშტაბს. ბევრი საბაზისო მოდელი ევროკავშირის ფარგლებს გარეთ მუშავდება და მათი ინტეგრაცია ადგილობრივ პროდუქტებსა თუ სერვისებში ხშირად ბუნდოვანს ხდის იურისდიქციულ საზღვრებს. “AI Act”-ის წარმატება დამოკიდებული იქნება ევროპელი მარეგულირებლების უნარზე, დაამკვიდრონ

ექსტრატერიტორიული გავლენა თანამშრომლობის მექანიზმების, ადეკვატურობის ჩარჩოებისა და საჯარო შესყიდვების წამახალისებელი მექანიზმების მეშვეობით, რომლებიც უპირატესობას ანიჭებენ შესაბამისობაში მყოფ სისტემებს.

შეჯამების სახით აღსანიშნავია, “AI Act”-ი გვთავაზობს უპრეცედენტო შესაძლებლობას, შეუთავსოს ტექნოლოგიური ინოვაცია დემოკრატიულ ღირებულებებსა და ფუნდამენტურ უფლებებს. თუმცა მისი იმპლემენტაცია მოითხოვს ძლიერ ინსტიტუტებს, სექტორთაშორის თანამშრომლობასა და მტკიცე პოლიტიკურ ნებას, რათა პასუხისმგებლიანი “AI” არა მხოლოდ მარეგულირებელ მისწრაფებად, არამედ ოპერაციულ რეალობად იქცეს.

8. დასკვნა: სამართლებრივი შესაბამისობიდან ეთიკურ და სტრატეგიულ მეურვეობამდე

მონაცემთა მინიმიზაციის პრინციპი, რომელიც ოდესღაც ტექნიკურ შეზღუდვად ან ბიუროკრატიულ ფორმალობად განიხილებოდა, ციფრულ ეპოქაში სათვალთვალ კაპიტალიზმის, ალგორითმული ექსპლუატაციისა და ძალაუფლების ასიმეტრიის წინააღმდეგ ნორმატიულ ბურჯად იქცა. ეპოქაში, რომელსაც სულ უფრო მეტად განსაზღვრავს მოდელის მაქსიმიზაცია და მონაცემთა მოდიფიკაცია, იგი ერთდროულად სამართლებრივ მოთხოვნასა და მორალურ იმპერატივს წარმოადგენს.

თუმცა ამ პრინციპის გამოყენება უნდა განვითარდეს იმ უნიკალური სირთულეების საპასუხოდ, რომლებსაც თანამედროვე “AI” სისტემები ქმნის. დიდი ენობრივი მოდელები და სხვა საბაზისო მოდელები ეჭვქვეშ აყენებს ტრადიციულ სამართლებრივ კატეგორიებსა და საპროცესო გარანტიებს, რაც მარეგულირებელი აღსრულებისადმი უფრო დინამიკურ და ჰოლისტიკურ მიდგომას მოითხოვს. შესაბამისად, მონაცემთა დაცვის საზედამხებელო ორგანოებმა (“DPA”) უნდა გაიაზრონ თავიანთი მანდატი — არა მხოლოდ ადასრულონ შესაბამისობა, არამედ ხელი შეუწყონ სისტემურ ანგარიშვალდებულებას, ეთიკურ რეფლექსიასა და საზოგადოებრივ ნდობას.

ეს გაფართოებული როლი გულისხმობს მონაცემთა გაუმართლებელი ან არაპროპორციული გამოყენების პრაქტიკისთვის წინააღმდეგობის გაწევას, გამჭვირვალე და განმარტებადი “AI”-ის ხელშეწყობასა და იმ ინდივიდების უფლებებისა და თავისუფლებების დაცვას, რომელთა მონაცემებიც ციფრული ინოვაციის საფუძველს წარმოადგენს. ასევე, ის მოითხოვს ინსტიტუციური შესაძლებლობების განვითარებას რისკზე დაფუძნებული აუდიტის ჩასატარებლად, სამოქალაქო საზოგადოებასთან თანამშრომლობისა და “AI” ტექნოლოგიების ეთიკურ მმართველობაში წვლილის შესატანად.

საბოლოო ჯამში, ხელოვნური ინტელექტის პასუხისმგებლიანი განვითარება და დანერგვა არ დაიყვანება სამართლებრივი ვალდებულებების ჩამონათვალზე. ეს არის კოლექტიური საზოგადოებრივი ვალდებულება,

რომლის მიზანია ტექნოლოგიური პროგრესის ბირთვში ადამიანის ღირსების, სამართლიანობისა და სამართლის პრინციპების დანერგვა. ამ ძალისხმევაში მონაცემთა დაცვის საზედამხებდველო ორგანოები ("DPA") არ არიან მხოლოდ მარეგულირებლები, ისინი ციფრული საზოგადოებრივი ინტერესის მეურვეები არიან.

ბიბლიოგრაფია:

1. CNIL, First Recommendations on the Development of AI Systems, Commission Nationale de l'Informatique et des Libertés, 2024, <<https://www.cnil.fr/en/ai-cnil-publishes-its-first-recommendations-development-artificial-intelligence-systems>> [23.07.2025].
2. Deng Z., Ma W., Han Q. L., Zhou W., Zhu X., Exploring DeepSeek: A Survey on Advances, Applications, Challenges and Future Directions, IEEE/CAA Journal of Automatica Sinica, 12(5), 2025, 872–893.
3. European Data Protection Board, Statement 3/2024 on the role of DPAs under the AI Act, 2024, <https://www.edpb.europa.eu/our-work-tools/our-documents/statements/statement-32024-data-protection-authorities-role-artificial_en> [23.07.2025].
4. European Data Protection Board, Opinion 28/2024 on Training Data for LLMs, 2024 <https://www.edpb.europa.eu/our-work-tools/our-documents/opinion-board-art-64/opinion-282024-certain-data-protection-aspects_en> [23.07.2025].
5. Garante per la protezione dei dati personali, Decision on OpenAI (ChatGPT), 2023, <[https://gdprhub.eu/index.php?title=Garante_per_la_protezione_dei_dati_personali_\(Italy\)_-_9870832](https://gdprhub.eu/index.php?title=Garante_per_la_protezione_dei_dati_personali_(Italy)_-_9870832)> [23.07.2025].
6. Hoofnagle C. J., van der Sloot B., Zuiderveen Borgesius F., The European Union General Data Protection Regulation: What it is and what it means. Information & Communications Technology Law, 28(1), 2019, 65–98.
7. ICO, The Lawful Basis for Web Scraping to Train Generative AI Models, Information Commissioner's Office (UK), 2024, <<https://ico.org.uk/about-the-ico/what-we-do/our-work-on-artificial-intelligence/response-to-the-consultation-series-on-generative-ai/the-lawful-basis-for-web-scraping-to-train-generative-ai-models/>> [23.07.2025].
8. Kuner C., Bygrave L. A., Docksey C., The EU General Data Protection Regulation (GDPR): A commentary. Oxford University Press, 2020.
9. Sangaraju V. V., AI and Data Privacy in Healthcare: Compliance with HIPAA, GDPR, and Emerging Regulations, International Journal of Emerging Trends in Computer Science and Information Technology, 2025, 67–74.
10. Vaezi A., Legal Challenges in the Deployment of Large Language Models: A Comparative Analysis under the GDPR and EU AI Act, 2025.