

საკვანძო სიტყვები: ხელოვნური ინტელექტი, კეთილსინდისიერება, სამეცნიერო კვლევა, ეთიკა.

შესავალი

ხელოვნური ინტელექტის გამოყენება სამეცნიერო კვლევაში ბევრ სარგებელს სთავაზობს სამეცნიერო საზოგადოებას, თუმცა ასევე ქმნის ახალ და რთულ ეთიკურ გამოწვევებს. მიუხედავად იმისა, რომ აღნიშნული ეთიკური საკითხები არ მოითხოვს მეცნიერების დადგენილი ეთიკური ნორმების შეცვლას, ისინი მოითხოვს სამეცნიერო საზოგადოების მხრიდან ხელოვნური ინტელექტის სათანადო გამოყენების ახალი სახელმძღვანელო პრინციპების შემუშავებას.

ხელოვნური ინტელექტის ინსტრუმენტები გამოიყენება სხვადასხვა სამეცნიერო ამოცანის შესასრულებლად, მათ შორის: მონაცემების ანალიზისთვის; მონაცემებისა და სურათების ინტერპრეტაციისთვის; ჰიპოთეზების გენერირებისთვის; კომპლექსური საკითხების მოდელირებისთვის; ჰიპოთეზებისა და მოდელების ვალიდაციისთვის; სამეცნიერო ლიტერატურის ძიება და მიმოხილვისთვის; სამეცნიერო ნაშრომებისა და სხვა კვლევითი შედეგების გაანალიზებისთვის. მიუხედავად იმისა, რომ ხელოვნური ინტელექტის გამოყენება სამეცნიერო კვლევაში გაიზარდა, ეთიკური პრინციპები მის გამოყენებასთან დაკავშირებით საჭიროებს განახლებას.

მაგალითად, კვლევის მთლიანობის ევროპული ქცევის კოდექსის 2023 წლის გადასინჯვა მოკლედ განიხილავს გამჭვირვალობის მნიშვნელობას. კოდექსი ითვალისწინებს, რომ მკვლევრებმა უნდა წარმოადგინონ ანგარიში მათი შედეგებისა და მეთოდების შესახებ, მათ შორის გარე სერვისების ან ხელოვნური ინტელექტისა და ავტომატიზირებული ინსტრუმენტების გამოყენების შესახებ და შინაარსის ფორმულირებისას ან პუბლიკაციების ფორმირებისას ხელოვნური ინტელექტის ან ავტომატიზირებული ინსტრუმენტების გამოყენების „დამალვა“ მიიჩნევა აკადემიური კეთილსინდისიერების პრინციპის დარღვევად.¹

ერთ-ერთი განახლებული ინსტიტუციური დოკუმენტი, „ჯანმრთელობის ეროვნული ინსტიტუტების სახელმძღვანელო პრინციპები და კვლევის ჩატარების პოლიტიკა“ იძლევა მითითებებს ხელოვნური ინტელექტის გამოყენების შესახებ.² ხელოვნური ინტელექტის შესახებ ეთიკის კოდექსები, როგორიცაა იუნესკოს „ხელოვნური ინტელექტის ეთიკა“³ და მეცნიერებისა და ტექნოლოგიების პოლიტიკის ოფისის „ხელოვნური ინტელექტის უფლებების ბილის გეგმა“,⁴ იძლევა სასარგებლო მითითებებს ზოგადად ხელოვნური ინტელექტის განვითარებისა და გამოყენებისთვის, თუმცა სამეცნიერო კვლევაში ხელოვნური ინტელექტის განვითარებისა და გამოყენების შესახებ კონკრეტული მითითებები არ არის მოცემული.

ამრიგად, სამეცნიერო კვლევაში ხელოვნური ინტელექტის გამოყენებასთან დაკავშირებით ეთიკურ დონეზე არსებობს ხარვეზი, რომელიც უნდა შეივსოს მისი სათანადო გამოყენების

¹ The European Code of Conduct for Research Integrity. 2023. Revised edition. <https://allea.org/code-of-conduct/>

² National Institutes of Health. 2023. *Guidelines for the Conduct of Research in the Intramural Program of the NIH.* https://oir.nih.gov/system/files/media/file/2023-11/guidelines-conduct_research.pdf

³ UNESCO. 2024. *Ethics of Artificial Intelligence.* <https://www.unesco.org/en/artificial-intelligence/recommendation-ethics>

⁴ Office of Science and Technology Policy. 2000. “Federal Research Misconduct Policy.” *Federal Register* 65 (235): 76260–76264.

ხელშეწყობის მიზნით. სამეცნიერო ნაშრომის ფორმირებისას ხელოვნური ინტელექტის გამოყენება წარმოშობს იმგვარ ეთიკურ საკითხებს, რომლებიც დაკავშირებულია ობიექტურობასთან, გამჭვირვალობასთან, კეთილსინდისიერებასთან, ანგარიშვალდებულებასთან, პასუხისმგებლობასა და მეცნიერებისადმი ნდობასთან.⁵ ხელოვნური ინტელექტის გამოყენება არ მოითხოვს მეცნიერების ეთიკური ნორმების რადიკალურ ცვლილებას, ის მოითხოვს სამეცნიერო საზოგადოებისგან ხელოვნური ინტელექტის სათანადო გამოყენების ახალი სახელმძღვანელო პრინციპების შემუშავებას. აღნიშნულიდან გამომდინარე, კეთილსინდისიერების პრინციპისა და ხელოვნური ინტელექტის ეთიკური გავლენის ანალიზი სამეცნიერო ნაშრომების ფორმირებაზე წარმოადგენს მეტად აქტუალურ საკითხს კვლევისათვის.

ძირითადი ნაწილი

მნიშვნელოვანია მეცნიერების ეთიკური ნორმების ფარგლებში გაანალიზდეს ხელოვნური ინტელექტის გამოყენება სამეცნიერო კვლევაში. მეცნიერების ეთიკური ნორმები არის პრინციპები, ღირებულებები, რომლებიც აუცილებელია სანდო, ვალიდური სამეცნიერო კვლევის ჩასატარებლად.⁶ აღნიშნული პრინციპების უმეტესობა განსაზღვრულია სამეცნიერო ქცევის კოდექსებში, პროფესიულ სახელმძღვანელო პრინციპებში, ინსტიტუციურ ან ჟურნალურ პოლიტიკაში, ან სამეცნიერო მეთოდოლოგიის შესახებ წიგნებსა და ნაშრომებში.⁷

სხვადასხვა სფეროს მეცნიერი, მათ შორის ფილოსოფიის, სოციოლოგიის, ისტორიის, სწავლობდა მეცნიერების ეთიკურ ნორმებს.⁸ სოციოლოგები, როგორებიც არიან მერტონი⁹ და შაპინი¹⁰ ეთიკურ ნორმებს განიხილავენ, როგორც განზოგადებებს, რომლებიც ზუსტად აღწერენ მეცნიერების პრაქტიკას, ხოლო ფილოსოფოსები, როგორებიც არიან კიტჩერი¹¹ და ჰააკი,¹² ამ ნორმებს განიხილავენ, როგორც წინასწარ განსაზღვრულ სტანდარტებს, რომლებიც მეცნიერებმა უნდა დაიცვან.

ხელოვნური ინტელექტი არის ქვესფერო კომპიუტერული მეცნიერების დისციპლინაში.¹³ თუმცა, ტერმინი „ხელოვნური ინტელექტი“ ასევე ხშირად გამოიყენება ტექნოლოგიების (ან ინსტრუმენტების) აღსანიშნავად, რომლებსაც შეუძლიათ შეასრულონ ადამიანური ამოცანები, რომლებიც მოითხოვს ინტელექტს, როგორიცაა აღქმა, განსჯა, მსჯელობა ან გადაწყვეტილების მიღება. მეცნიერების აზრით, არცერთი ხელოვნური ინტელექტის სისტემა არ არის შეცდომებისგან თავისუფალი. მნიშვნელოვანია „შეცდომის“ განმარტება და „სისტემური შეცდომებისა“ და „შემთხვევითი შეცდომების“ ერთმანეთისგან გარჩევა. სიტყვა „შეცდომა“

⁵ Alvarado, R. 2023. “AI as an Epistemic Technology.” *Science and Engineering Ethics* 29: 32. <https://link.springer.com/article/10.1007/s11948-023-00451-3>

⁶ Resnik, D.B., and K.C. Elliott. 2019. “Value-Entanglement and the Integrity of Scientific Research.” *Stud. Hist. Philos. Sci.* 75: 1–11.

⁷ International Committee of Medical Journal Editors. 2023. *Recommendations for the Conduct, Reporting, Editing, and Publication of Scholarly Work in Medical Journals*. <https://www.icmje.org/icmje-recommendations.pdf>

⁸ Howson, C., and P. Urbach. 2005. *Scientific Reasoning: A Bayesian Approach*. 3rd edn. Open Court, New York.

⁹ Merton, R. 1973. *The Sociology of Science*. University of Chicago Press, Chicago.

¹⁰ Shapin, S. 1995. “Here and Everywhere: Sociology of Scientific Knowledge.” *Ann. Rev. Sociol.* 21: 289–321.

¹¹ Kitcher, P. 1993. *The Advancement of Knowledge*. Oxford University Press, New York.

¹² Haack, S. 2007. *Defending Science within Reason*. Prometheus Books, New York.

¹³ Licht, K., and J. Licht. 2020. “Artificial Intelligence, Transparency, and Public Decision-Making: Why Explanations Are Key When Trying to Produce Perceived Legitimacy.” *AI Soc.* 35: 917–926.

სხვადასხვა მნიშვნელობა აქვს, აქ იგულისხმება გრამატიკული შეცდომები, მსჯელობიდან გამომდინარე შეცდომები, ტიპოგრაფიულ შეცდომები, გაზომვის შეცდომები და ა.შ.¹⁴ „შეცდომის“ ამ სხვადასხვა მნიშვნელობას საერთო აქვს ის, რომ: (1) შეცდომები გულისხმობს სისწორის სტანდარტიდან გადახრას; და (2) შეცდომები, როდესაც ისინი ცნობიერი არსებების მიერ არის დაშვებული, უნებლიეა; ანუ ისინი შემთხვევითობა ან შეცდომაა და განსხვავდება თაღლითობისა და მოტყუებისგან.¹⁵

სისტემურ და შემთხვევით შეცდომებს შორის განსხვავება შეიძლება ბუნდოვანი იყოს, რადგან შემთხვევითი სახის შეცდომები შეიძლება სისტემურ ხასიათს ატარებდეს.¹⁶ მიუხედავად ამისა, სისტემური შეცდომები ხშირად უფრო საზიანოა მეცნიერებისა და საზოგადოებისთვის, ვიდრე შემთხვევითი შეცდომები.

ხელოვნური ინტელექტის სისტემებს შეუძლიათ შემთხვევითი შეცდომების დაშვება.¹⁷ ეს პრობლემა აშკარაა განსაკუთრებით ბიზნესის, სამართლისა და სამეცნიერო კვლევების სფეროში. მაგალითად, „ChatGPT“ მიდრეკილია შემთხვევითი ფაქტობრივი ხასიათის და ციტირებისას შეცდომების დაშვებისკენ. მაგალითად, მკვლევარებმა ბჰატაჩარიამ და სხვებმა თავიანთი სამეცნიერო კვლევის ფარგლებში გამოიყენეს „ChatGPT 3.5“ სამედიცინო თემებზე 30 მოკლე ნაშრომის (გენერირებისთვის. ჩატბოტის მიერ წარმოებული ინფორმაციის 47%-ი იყო შეთხზული, 46%-ი იყო ავთენტური, მაგრამ მცდარად გამოყენებული და მხოლოდ 7%-ი იყო ჭეშმარიტი.¹⁸ მიუხედავად იმისა, რომ „ChatGPT 4.0“ უკეთ მუშაობს, ვიდრე „ChatGPT 3.5“, ის მაინც იყენებს შეთხზულ და მცდარ ციტატებს.¹⁹ აღსანიშნავია შემთხვევა, როდესაც აშშ-ში ორმა ამერიკელმა ადვოკატმა ChatGPT-ის მიერ მომზადებული სასამართლო დოკუმენტაცია წარადგინა სასამართლოში, რომელიც შეიცავდა არაზუსტ ინფორმაციას, ციტატებს, რის გამოც მოსამართლემ ისინი 5000 (ხუთი ათასი) დოლარით დააჯარიმა. მოსამართლემ განაცხადა, რომ იგი არ იყო ChatGPT-ის გამოყენების მოწინააღმდეგე, თუმცა ადვოკატებს სათანადო, გულისხმიერი ყურადღება უნდა გამოეჩინათ ხელოვნური ინტელექტის მიერ გენერირებული ტექსტის სიზუსტის გადამოწმებისას.²⁰

სამეცნიერო ლიტერატურაში განხილული გენერაციული ხელოვნური ინტელექტის მიერ დაშვებული შემთხვევითი შეცდომების მაგალითია ყალბი ციტატების შემთხვევა. ერთ-ერთი მიზეზი, რის გამოც მაგალითად, „ChatGPT“-ი, ყალბ, მაგრამ ერთი შეხედვით რეალურ ციტატებს იყენებს, არის ის, რომ პროგრამა ტექსტურ მონაცემებს ადამიანებისგან განსხვავებულად ამუშავებს.²¹ მკვლევრები ციტატებს იყენებენ კონკრეტული ტექსტის წაკითხვიდან გამომდინარე და მისი ციტირებით, მაგრამ „ChatGPT“-ი ციტატებს ქმნის ტექსტური მონაცემების დიდი

¹⁴ იხ. იქვე.

¹⁵ Mitchell, M. 2019. *Artificial Intelligence*. Picador, New York.

¹⁶ Alvarado, R. 2022. „Computer Simulations as Scientific Instruments.” *Found. Sci.* 27(3): 1183–1205.

¹⁷ Mitchell, M. 2019. *Artificial Intelligence*. Picador, New York.

¹⁸ Bhattacharyya, M., V.M. Miller, D. Bhattacharyya, and L.E. Miller. 2023. „High Rates of Fabricated and Inaccurate References in ChatGPT-Generated Medical Content.” *Cureus* 15(5): e39238.

¹⁹ Walters, W.H., and E.I. Wilder. 2023. „Fabrication and Errors in the Bibliographic Citations Generated by ChatGPT.” *Sci. Rep.* 13: 14045.

²⁰ Milmo, D. 2023. „Two US Lawyers Fined for Submitting Fake Court Citations from ChatGPT.” *The Guardian*. <https://www.theguardian.com/technology/2023/jun/23/two-us-lawyers-fined-submitting-fake-court-citations-chatgpt>

²¹ Walters, W.H., and E.I. Wilder. 2023. „Fabrication and Errors in the Bibliographic Citations Generated by ChatGPT.” *Sci. Rep.* 13: 14045.

რაოდენობის დამუშავების გზით და ციტირების მოთხოვნაზე მაღალი ალბათობით პასუხის გენერირებით.²²

გარდა ზემოხსენებულისა, კვლევებმა აჩვენა, რომ რასობრივი და ეთნიკური მიკერძოებები გავლენას ახდენს ხელოვნური ინტელექტის/მანქანური სწავლების გამოყენებაზე სამედიცინო ვიზუალიზაციაში, დიაგნოზსა და პროგნოზირებაში ჯანდაცვის მონაცემთა ბაზებში არსებული მიკერძოებების გამო.²³ მეცნიერები წლების განმავლობაში ებრძოდნენ კვლევაში მიკერძოებას, შესაბამისად, შეიმუშავეს მეთოდები და სტრატეგიები ექსპერიმენტული დიზაინის, მონაცემთა ანალიზის, მოდელებისა და თეორიის აგების მიკერძოების მინიმიზაციისა და კონტროლისთვის.²⁴ თუმცა, მეცნიერებაში ხელოვნური ინტელექტის გამოყენებასთან დაკავშირებული მიკერძოება შეიძლება იყოს რთულად აღმოსაჩენი კვლევითი მონაცემების ზომისა და სირთულის, ასევე მონაცემებს, ალგორითმებსა და აპლიკაციებს შორის ურთიერთქმედების გამო.²⁵ მეცნიერებმა, რომლებიც კვლევაში ხელოვნური ინტელექტის სისტემებს იყენებენ, უნდა გადადგან შესაბამისი ნაბიჯები მიკერძოებების პროგნოზირების, იდენტიფიცირების, კონტროლისა და მინიმიზაციის მიზნით, იმის უზრუნველყოფით, რომ გაამჟღავნონ მონაცემთა ანალიზში გამოყენებული ცვლადები, ალგორითმები, მოდელები და პარამეტრები.²⁶

შემთხვევითმა შეცდომებმა შეიძლება დიდი საფრთხე შეუქმნას სამეცნიერო ცოდნის ვალიდურობასა და სანდოობას.²⁷ მიუხედავად იმისა, რომ მეცნიერებაში ზოგიერთი შემთხვევითი შეცდომა გარდაუვალია, ხელოვნური ინტელექტის გამოყენებისას მისი გადაჭარბებული რაოდენობა შეიძლება ჩაითვალოს დაუდევრობად ან უგუნურებად. შემთხვევითი შეცდომების შემცირება ფართოდ არის აღიარებული, როგორც კარგი სამეცნიერო მეთოდოლოგიისა და პრაქტიკის აუცილებელი პირობა.²⁸ მიუხედავად იმისა, რომ კვლევაში ზოგიერთი შემთხვევითი შეცდომა გარდაუვალია, მეცნიერებს აქვთ ვალდებულება, ამოიცნონ, შეამცირონ და გამოსაწირონ ისინი, რადგან საბოლოო ჯამში ისინი პასუხისმგებელნი არიან როგორც ადამიანურ, ასევე ხელოვნური ინტელექტის მიერ გენერირებულ შეცდომებზე. აღნიშნულთან დაკავშირებით გამჭვირვალობის კრიტერიუმი მნიშვნელოვანია კვლევის სანდოობისა და ვალიდურობისთვის.²⁹

მეცნიერებაში შეცდომების შემცირების სტრატეგიები მოიცავს ხარისხის გაუმჯობესების ტექნიკას, როგორიცაა მონაცემების, ინსტრუმენტებისა და სისტემების აუდიტი; ინსტრუმენტების ვალიდაცია და ტესტირება, რომლებიც აანალიზებენ ან ამუშავებენ მონაცემებს; და შეცდომების

²² Bhattacharyya, M., V.M. Miller, D. Bhattacharyya, and L.E. Miller. 2023. "High Rates of Fabricated and Inaccurate References in ChatGPT-Generated Medical Content." *Cureus* 15(5): e39238.

²³ Garin, S.P., V.S. Parekh, J. Sulam, and P.H. Yi. 2023. "Medical Imaging Data Science Competitions Should Report Dataset Demographics and Evaluate for Bias." *Nat. Med.* 29(5): 1038–1039.

²⁴ Giere, R., J. Bickle, and R.F. Maudlin. 2005. *Understanding Scientific Reasoning*. 5th edn. Wadsworth, Belmont.

²⁵ Licht, K., and J. Licht. "Artificial Intelligence, Transparency, and Public Decision-Making: Why Explanations Are Key When Trying to Produce Perceived Legitimacy." *AI Soc.* 35: 917–926.

²⁶ Daneshjou, R., M.P. Smith, M.D. Sun, V. Rotemberg, and J. Zou. 2021. "Lack of Transparency and Potential Bias in Artificial Intelligence Data Sets and Algorithms: A Scoping Review." *JAMA Dermatol.* 157(11): 1362–1369.

²⁷ Shamoo, A.E., and D.B. Resnik. 2022. *Responsible Conduct of Research*. 4th edn. Oxford University Press, New York.

²⁸ იბ. იქვე.

²⁹ Ball, P. 2023. "Is AI Leading to a Reproducibility Crisis in Science?" *Nature* 624: 22–25.

გამოკვლევა და ანალიზი. ჟურნალის რეცენზია და გამოქვეყნების შემდგომი რეცენზია ასევე მნიშვნელოვან როლს ასრულებს შეცდომების შემცირებაში. ნებისმიერ შემთხვევაში, სამეცნიერო ანგარიშვალდებულება მოითხოვს, რომ მკვლევარებმა აიღონ პასუხისმგებლობა ხელოვნური ინტელექტის/მანქანური სწავლების სისტემების მიერ დაშვებულ შეცდომებზე, შეძლონ ახსნან შეცდომების წარმოშობის მიზეზები.

ხელოვნურ ინტელექტის მიერ გენერირებულ ტექსტში არსებულმა შეცდომებმა, შეიძლება გამოიწვიოს მკვლევრის პასუხისმგებლობა, თუ ისინი განზრახ, შეგნებულად ან დაუდევრად ავრცელებენ ცრუ მონაცემებს ან ადგილი ექნება პლაგიატს. მიუხედავად იმისა, რომ დანაშაულის შესახებ რეგულაციებისა და პოლიტიკის უმეტესობა განასხვავებს დარღვევასა და უნებლიე შეცდომას, მეცნიერები მაინც შეიძლება იყვნენ პასუხისმგებელნი უნებლიედ დაშვებულ შეცდომაზე.³⁰ მაგალითად, პირი, რომელიც იყენებს ChatGPT-ს სამეცნიერო ნაშრომის დასაწერად შეცდომების ან პლაგიატის შედეგების შემოწმების გარეშე, შეიძლება პასუხისმგებელი იყოს კვლევაში დაშვებული უზუსტობებისთვის ხელოვნური ინტელექტის დაუდევრად გამოყენებისთვის. შესაბამისად, ხელოვნური ინტელექტის გამოყენება მოითხოვს შესაბამისი მკვლევრის მიერ ადეკვატური ნაბიჯების გადადგმას შეცდომების მინიმიზაციისა და კონტროლისთვის.

ხელოვნური ინტელექტის გამოყენება კვლევაში, წარმოშობს ეთიკურ საკითხებს მონაცემთა კონფიდენციალურობასთან დაკავშირებით. მაგალითად, “ChatGPT”-ი ინახავს მომხმარებლების მიერ გაზიარებულ ინფორმაციას, მათ შორის მონაცემებს, რომლებიც მკვლევრის მიერ წარდგენილია საწყისი მოთხოვნებისა და შემდგომი კომუნიკაციის დროს.

ასევე, სამეცნიერო ნაშრომების ავტორობა არა მხოლოდ მასში მნიშვნელოვანი წვლილის შეტანაზეა დაფუძნებული, არამედ ნაშრომზე პასუხისმგებლობის აღებაზეც. პირი არ უნდა იყოს დასახელებული ნაშრომის ავტორი, თუ მას არ შეუძლია პასუხისმგებლობის აღება ნაშრომში შეტანილი წვლილისთვის. თუ გამოქვეყნების შემდეგ ნაშრომთან დაკავშირებით კითხვები წარმოიშვება, ავტორს უნდა შეეძლოს ამ კითხვებზე გასაგებად პასუხის გაცემა. ხელოვნური ინტელექტის სისტემები არ შეიძლება პასუხისმგებელნი იყვნენ თავიანთ ქმედებებზე შემდეგი ორი მიზეზიდან გამომდინარე: 1) ხელოვნურ ინტელექტს არ შეუძლია შესაბამისი, ვალიდური ახსნა-განმარტებების წარმოდგენა გენერირებული ინფორმაციის შესახებ, 2) ხელოვნურ ინტელექტს არ შეუძლია მორალურად პასუხისმგებელი იყოს მის ქმედებებზე, 3) ხელოვნურ ინტელექტს ვერ დაეკისრება მორალური, სამართლებრივი, თუ ფინანსური სანქცია. მიუხედავად იმისა, რომ ხელოვნური ინტელექტის სისტემებს ვერ დაეკისრება პასუხისმგებლობა, მაინც მნიშვნელოვანია სამეცნიერო ნაშრომში მათი წვლილის სათანადოდ აღიარება. აღიარება საჭიროა კვლევაში კეთილსინდისიერებისა და გამჭვირვალობის ხელშეწყობის მიზნით.

მკვლევრები დევიდ რეზნიკი და მოჰამად ჰოსეინი თავიანთ სამეცნიერო კვლევაში “The ethics of using artificial intelligence in scientific research: new guidance needed for a new tool”³¹ წარმოადგენენ

³⁰ Caron, M.M., S.B. Dohan, M. Barnes, and B.E. Bierer. 2023. “Defining ‘Recklessness’ in Research Misconduct Proceedings.” *Accountability in Research*: 1–23.

³¹ ხელოვნური ინტელექტის სამეცნიერო კვლევაში გამოყენების ეთიკა: ახალი ინსტრუმენტისთვის ახალი სახელმძღვანელოს საჭიროება (აქ და ყველგან მთარგმნ. ქრისტინე უბილავა).

რეკომენდაციებს ხელოვნური ინტელექტის სამეცნიერო კვლევაში გამოყენების ეთიკასთან დაკავშირებით, რომლებიც შემდეგნაირად გამოიყურება:

- მკვლევრები პასუხისმგებელნი არიან ხელოვნურ ინტელექტის მიერ გენერირებულ ტექსტში შემთხვევითი შეცდომების იდენტიფიცირებაზე, შემცირებასა და კონტროლზე.
- მკვლევრებმა უნდა ახსნან ხელოვნური ინტელექტის გამოყენების აუცილებლობა კვლევაში.
- მკვლევრებმა ხელოვნური ინტელექტი უნდა გამოიყენონ მხოლოდ მაშინ, როდესაც მათ გააჩნია საკმარისი ცოდნა დასმულ საკითხთან დაკავშირებით.
- მკვლევარები, რომლებიც განზრახ, შეგნებულად ან დაუფიქრებლად იყენებენ ხელოვნურ ინტელექტს სამეცნიერო ნაშრომის შესაქმნელად ან გასაყალბებლად, პასუხისმგებელნი არიან თავიანთ ქცევაზე.
- სინთეზური მონაცემების გამოყენებით მკვლევრებმა უნდა: (1) მიუთითონ მონაცემების, თუ რომელი ნაწილებია სინთეზური; (2) ნათლად მონიშნონ სინთეზური მონაცემები; (3) აღწერონ, თუ როგორ იქნა გენერირებული მონაცემები; და (4) ახსნან, თუ როგორ და რატომ იქნა გამოყენებული მონაცემები.³²

ჟურნალებმა და ინსტიტუტებმა უნდა განსაზღვრონ მკაფიო და აღსრულებადი სახელმძღვანელო პრინციპები. „პუბლიკაციების ეთიკის კომიტეტის“ (COPE)³³ ინიციატივით დადგინდა ერთი მხრივ ის, რომ ხელოვნური ინტელექტის სისტემები ავტორებად ვერ იქნებიან მითითებული, რადგან მათ არ შეუძლიათ აკადემიური პასუხისმგებლობის აღება და მეორე მხრივ, „პუბლიკაციების ეთიკის კომიტეტი“ მოითხოვს, რომ სამეცნიერო კვლევის პროცესში გამოყენებული ხელოვნური ინტელექტის ინსტრუმენტები გამჭვირვალედ იყოს მითითებული ნაშრომში. იუნესკოს რეკომენდაცია ხელოვნური ინტელექტის ეთიკის შესახებ ასევე ხაზს უსვამს ხელოვნურ ინტელექტზე დაფუძნებულ კვლევაში ადამიანური გადამოწმებადობის, ზედამხედველობის მნიშვნელობაზე.

ხელოვნურ ინტელექტთან დაკავშირებული რისკების მართვისთვის ასევე აუცილებელია ტექნოლოგიური ზომების მიღება. ხელოვნური ინტელექტის მიერ გენერირებული კონტენტისთვის შესაბამისი ნიშნის სისტემები, რომლებიც შემუშავებულია ისეთი ორგანიზაციების მიერ, როგორიცაა “OpenAI”-ი და “DeepMind”-ი, მიზნად ისახავს მანქანური გენერირებული ტექსტის ადამიანის მიერ შექმნილი შინაარსისგან გარჩევას. ანალოგიურად, „ახსნილი ხელოვნური ინტელექტის“ (XAI)³⁴ განვითარება თანდათან უწყობს ხელს

³² Resnik, D.B., and M. Hosseini. 2025. “The Ethics of Using Artificial Intelligence in Scientific Research: New Guidance Needed for a New Tool.” *AI Ethics* 5(2): 1499–1521. doi:10.1007/s43681-024-00493-8.

³³ The Committee on Publication Ethics - არის საერთაშორისო არაკომერციული ორგანიზაცია, რომელიც მუშაობს სამეცნიერო პუბლიკაციების ეთიკისა და აკადემიური კეთილსინდისიერების მიმართულებით. ორგანიზაცია დაარსდა 1997 წელს დიდ ბრიტანეთში სამედიცინო ჟურნალების რედაქტორების მიერ. მისი მთავარი მიზანია სამეცნიერო გამოცემებში ეთიკური სტანდარტების დამკვიდრება და გამომცემლების, რედაქტორებისა და ავტორების მხარდაჭერა.

³⁴ Explainable Artificial Intelligence არის მიმართულება, რომლის მიზანია ხელოვნური ინტელექტის სისტემების გადაწყვეტილებების ადამიანისთვის გასაგებად ახსნა.

გამჭვირვალობას, რაც ხელოვნურ ინტელექტთან დაკავშირებული გადაწყვეტილების მიღების პროცესს უფრო გასაგებს და ხელმისაწვდომს ხდის მკვლევრებისთვის.³⁵

საგანმანათლებლო ინიციატივები, როგორიცაა „ჯანმრთელობის ეროვნული ინსტიტუტების“ (NIH)³⁶ და უნივერსიტეტების მიერ წარმართული ხელოვნური ინტელექტის ეთიკის პროგრამები ხელს უწყობს მეცნიერებში პასუხისმგებლობის გრძნობისა და ეთიკური ცნობიერების ამაღლებას. მაგალითად, უნივერსიტეტების მიერ შექმნილი კურსები ხელოვნური ინტელექტის ეთიკის შესახებ უზიარებს მეცნიერებს იმ ცოდნასა და ინსტრუმენტებს, რომლებიც აუცილებელია ხელოვნური ინტელექტის გამოყენებისათვის. გარდა ამისა, ჟურნალების რედაქტორებს შორის ინტერდისციპლინარული თანამშრომლობა აუცილებელია.³⁷

მკვლევრებს შეუძლიათ გამოიყენონ ხელოვნური ინტელექტის ინსტრუმენტები კონკრეტული, შეზღუდული მიზნებისთვის. თუმცა, ხელოვნური ინტელექტი ვერ შეცვლის ავტორების ინტელექტუალურ წვლილს ან მათი ნაშრომის მიმართ პასუხისმგებლობას. გამჭვირვალობის უზრუნველსაყოფად, ავტორებს მოეთხოვებათ გაამჟღავნონ ხელოვნური ინტელექტის ინსტრუმენტების ნებისმიერი გამოყენება სამეცნიერო ნაშრომის მომზადებისას. თუ ხელოვნური ინტელექტი გამოიყენება ენის რედაქტირებისთვის, გრამატიკული ნორმების შემოწმებისა ან ლიტერატურის მოძიებისთვის, ეს ნათლად უნდა იყოს მითითებული სამადლობელო გვერდზე ან სქოლიოში. თუმცა, მაგალითად, საერთაშორისო აკადემიური გამომცემლობა “Taylor & Francis”-ის³⁸ მოქმედი ხელოვნური ინტელექტის პოლიტიკის თანახმად, ხელოვნური ინტელექტის მიერ გენერირებული სურათები, მათ შორის ფიგურები, დიაგრამები და მონაცემთა ცხრილები, არ არის დაშვებული სამეცნიერო ნაშრომში გამოსაყენებლად.

ხელოვნურმა ინტელექტმა არ უნდა ჩაანაცვლოს ადამიანის სამეცნიერო ნაშრომი. ხელოვნური ინტელექტი არ შეიძლება ჩაითვალოს ავტორად. ავტორობა მოითხოვს ინტელექტუალურ წვლილს, ანგარიშვალდებულებას და პასუხისმგებლობას, რომლის შესრულებაც ხელოვნური ინტელექტის ინსტრუმენტებს არ შეუძლიათ. მკვლევრებს ევალებათ: 1) თავიდან აიცილონ ხელოვნურ ინტელექტზე სრული ან ნაწილობრივი დამოკიდებულება ლიტერატურული მიმოხილვების, ჰიპოთეზების, დისკუსიებისა ან დასკვნების გენერირებისთვის, 2) უზრუნველყონ თავიანთი ნაშრომის ორიგინალურობა და დაადასტურონ, რომ ხელოვნური ინტელექტის მიერ გენერირებული ტექსტი არ ეწინააღმდეგება კეთილსინდისიერების პრინციპს ან არ იწვევს პლაგიატს, 3) წარმოადგინონ კვლევაში გამოყენებული ხელოვნური ინტელექტის ინსტრუმენტებისა და მოდელების მკაფიო ახსნა-განმარტებები.

შირი განდი და სხვები თავიანთ სამეცნიერო სტატიაში “AI’s Impact on Academic Research Integrity: Guidelines from the Editors of Journal of Global Marketing”³⁹ აღნიშნავენ, რომ ხელოვნური

³⁵ Gunning, D., and D. Aha. 2019. “DARPA’s Explainable Artificial Intelligence (XAI) Program.” *AI Magazine* 40(2): 44–58. <https://doi.org/10.1609/aimag.v40i2.2850>

³⁶ National Institutes of Health.

³⁷ Floridi, L., and J. Cowls. 2022. “A Unified Framework of Five Principles for AI in Society.” In *Machine Learning and the City: Applications in Architecture and Urban Design*, edited by Silvio Carta. John Wiley & Sons Online Library.

³⁸ იგი დაარსდა 1852 წელს დიდ ბრიტანეთში და დღეს ერთ-ერთ ყველაზე მსხვილ აკადემიურ გამომცემლად ითვლება მსოფლიოში.

³⁹ ხელოვნური ინტელექტის გავლენა აკადემიური კვლევის მთლიანობაზე: გლობალური მარკეტინგის ჟურნალის რედაქტორების რეკომენდაციები.

ინტელექტის გამოყენებისას მკვლევარებს ევალებათ ის, რომ მათი ნამუშევარი შეესაბამებოდეს დადგენილ ეთიკურ სტანდარტებს, მათ შორის „პუბლიკაციების ეთიკის კომიტეტის“ (COPE) მიერ დადგენილ სტანდარტებს. კერძოდ:

- ხელოვნური ინტელექტის საშუალებით გენერირებული ტექსტის ფარგლებში მონაცემთა კონფიდენციალურობის შესაბამისი რეგულაციების, როგორცაა GDPR-ის⁴⁰ ან CCPA-ს,⁴¹ დაცვა.
- ხელოვნური ინტელექტი არ უნდა იქნას გამოყენებული მონაცემების ან კვლევის შედეგების გაყალბების მიზნით.
- მკვლევარებმა ცალსახად უნდა მიუთითონ ნაშრომში ხელოვნური ინტელექტის მიერ გენერირებული ტექსტის შესახებ, რათა უზრუნველყონ სიზუსტე, ორიგინალურობა და ეთიკურ სამეცნიერო პრაქტიკასთან შესაბამისობა.
- წარდგენილი ნაშრომის ორიგინალურობაზე საბოლოო პასუხისმგებლობა ეკისრებათ ადამიან ავტორებს.

ამავდროულად, რედაქტორებმა და რეცენზენტებმა არ უნდა გამოიყენონ ხელოვნური ინტელექტის ინსტრუმენტები რეცენზირების ანგარიშების მოსამზადებლად. რეცენზენტებმა უნდა დაიცვან მკაცრი კონფიდენციალურობა და თავი შეიკავონ ნამუშევრის შინაარსის ხელოვნური ინტელექტის ინსტრუმენტებში ატვირთვისგან.⁴²

ხელოვნური ინტელექტის მიერ გენერირებული სამეცნიერო ნაშრომების რაოდენობის ზრდის გამო, მკვლევარები ავითარებენ ალგორითმებს მათი ამოცნობის გასაუმჯობესებლად.⁴³ სამეცნიერო ჟურნალებმა დაიწყეს ავტორების პოლიტიკის სახელმძღვანელო პრინციპების შემუშავება. მიუხედავად იმისა, რომ გამომცემლობები ეთანხმებიან ხელოვნური ინტელექტის გამოყენებას, ეთიკური პრინციპების დაცვა აღნიშნულ პროცესში აუცილებელია. ასევე, საჭიროა ხელოვნური ინტელექტის სწავლებასა და კვლევაში გამოყენების, ასევე მისი გამოყენების პასუხისმგებლობებისა და შედეგების შემდგომი კვლევების ჩატარება.

მიუხედავად იმისა, რომ ხელოვნური ინტელექტის ინსტრუმენტები არაერთ სარგებელს სთავაზობს მკვლევარს, მათმა არაკეთილსინდისიერად გამოყენებამ შეიძლება უარყოფითად იმოქმედოს კვლევის რეფლექსურ და თვითრეფლექსურ ბუნებაზე. თუმცა, მისი ეფექტურად, აკადემიური კეთილსინდისიერებისა და ეთიკური პრინციპების სწორად გამოყენების შემთხვევაში, ხელოვნურ ინტელექტს აქვს პოტენციალი, მნიშვნელოვნად დააჩქაროს სამეცნიერო ნაშრომების შექმნის პროცესი. ხელოვნური ინტელექტის სისტემებისა და ავტომატიზაციის ინსტრუმენტების სამეცნიერო ნაშრომებში ინტეგრირების მიზანი არ უნდა იყოს სამეცნიერო

⁴⁰ General Data Protection Regulation - არის ევროკავშირის მონაცემთა დაცვის რეგულაცია, რომელიც 2018 წელს შევიდა ძალაში.

⁴¹ California Consumer Privacy Act - არის კალიფორნიის შტატის მონაცემთა დაცვისა და მომხმარებელთა კონფიდენციალურობის კანონი, რომელიც ძალაში 2020 წელს შევიდა.

⁴² Gandhi, S., M.S. Ahmed, L. Huang, A. Behl, and A.K. Manrai. 2025. "AI's Impact on Academic Research Integrity: Guidelines from the Editors of Journal of Global Marketing." *Journal of Global Marketing* 38(3): 189–192. <https://doi.org/10.1080/08911762.2025.2488060>

⁴³ Deanda, D., I. Alsmadi, J. Guerrero, and G. Liang. 2025. "Defending Mutation-Based Adversarial Text Perturbation: A Black-Box Approach." *Clust. Comput.* 28: 196.

ნაშრომების მასობრივი წარმოება, არამედ მკვლევართა წერითი და ანალიტიკური უნარების გაუმჯობესება.

მკვლევრები, ჯო ლერენა-იზკიერდო და რაკელ აიალა-კარაბახო თავიანთ სამეცნიერო სტატიაში “Ethics of the Use of Artificial Intelligence in Academia and Research: The Most Relevant Approaches, Challenges and Topics”⁴⁴ აღნიშნავენ, რომ ხელოვნური ინტელექტის ინტეგრაცია სხვადასხვა აკადემიურ დისციპლინაში წარმოადგენს შესაძლებლობებისა და რისკების ფართო სპექტრს, რაც მოითხოვს მუდმივ ეთიკურ განსჯას.⁴⁵ ერთი მხრივ, ხელოვნური ინტელექტი მოქმედებს როგორც ძლიერი კატალიზატორი, რომელსაც შეუძლია დაამუშაოს დიდი რაოდენობის ინფორმაცია, რომელიც აღემატება ადამიანურ შესაძლებლობებს, ხელი შეუწყოს ინტერდისციპლინარულობას და კვლევითი რთული პროცესების მართვას. განათლებაში იგი ტრანსფორმაციულ პოტენციალს იძლევა პერსონალიზებული სწავლებისა და რესურსებზე წვდომის სახით. ხელოვნურ ინტელექტსა და ადამიანებს შორის თანამშრომლობა ფუნდამენტური პრინციპია, რომელიც უნდა იყოს დაცული სამეცნიერო ნაშრომზე მუშაობის პროცესში. ხელოვნურ ინტელექტს შეუძლია რუტინული ამოცანების ავტომატიზაცია, თუმცა სამეცნიერო ნაშრომში ანალიზი უნდა ეფუძნებოდეს ეთიკურ პრინციპებს, მკვლევრის ცოდნას, შედეგების შესახებ ცნობიერებას და მორალურ ღირებულებებს.

აუცილებელია, რომ ხელოვნური ინტელექტი აღიქმებოდეს, როგორც ინსტრუმენტი და არა როგორც ინტელექტი, რომლის აზრი ექვგარეშედ უნდა გამოიყენოს მკვლევარმა სამეცნიერო ნაშრომში. ხელოვნური ინტელექტის მიერ გენერირებული ტექსტის სამეცნიერო ნაშრომში გამოყენების ძირითადი პრინციპი უნდა იყოს კეთილსინდისიერება. ამ მიზნით, უნივერსიტეტებმა უნდა შეუწყონ ხელი გამჭვირვალობასა და ანგარიშვალდებულებას.

ერთ-ერთი მნიშვნელოვანი საკითხი ზემოხსენებულ პროცესში არის მითითების/„პრომპტის“ (Prompt) სწორად გამოყენება. „პრომპტის“ ფორმულირება პირდაპირ გავლენას ახდენს მიღებული ინფორმაციის სიზუსტეზე, სანდოობასა და აკადემიურ ხარისხზე. არაზუსტად ან ბუნდოვნად ფორმულირებულმა „პრომპტმა“ შეიძლება გამოიწვიოს მცდარი, არასრული ან ეთიკურად პრობლემური ინფორმაციის არსებობა, რაც საბოლოოდ საფრთხეს უქმნის სამეცნიერო ნაშრომის სანდოობას.

„პრომპტების“ ზუსტი ფორმულირება განსაკუთრებით მნიშვნელოვანია სხვადასხვა მიზეზიდან გამომდინარე. პირველ რიგში, სწორად ჩამოყალიბებული მოთხოვნა ხელს უწყობს სანდო და რელევანტური ინფორმაციის მიღებას, რაც ამცირებს დეზინფორმაციისა და კვლევაში შეცდომების დაშვების რისკს. მეორე მხრივ, „პრომპტის“ ფორმულირება განსაზღვრავს, რამდენად გამჭვირვალე იქნება ხელოვნური ინტელექტის გამოყენება კვლევის პროცესში. მაგალითად, თუ მკვლევარი ხელოვნურ ინტელექტს პირდაპირ მოსთხოვს დამოუკიდებლად ჩამოაყალიბოს

⁴⁴ ხელოვნური ინტელექტის გამოყენების ეთიკა აკადემიურ წრეებსა და კვლევაში: ყველაზე აქტუალური მიდგომები, გამოწვევები და თემები.

⁴⁵ Llerena-Izquierdo, J., and R. Ayala-Carabajo. 2025. “Ethics of the Use of Artificial Intelligence in Academia and Research: The Most Relevant Approaches, Challenges and Topics.” *Informatics* 12(4): 111. <https://doi.org/10.3390/informatics12040111>.

არგუმენტები წყაროების გარეშე, იზრდება პლაგიატის, ფაქტობრივი შეცდომებისა და არაკეთილსინდისიერების რისკი.

აღნიშნულ კონტექსტში განსაკუთრებული მნიშვნელობა ენიჭება იმასაც, რომ მკვლევარმა მკაფიოდ განსაზღვროს ხელოვნური ინტელექტის როლი კვლევის პროცესში, კერძოდ, გამოიყენება, თუ არა იგი მხოლოდ ტექსტის რედაქტირებისთვის, იდეების გენერირებისა და მონაცემთა ანალიზისთვის. გარდა ამისა, „პრომპტების“ სწორ ფორმულირებას აქვს არა მხოლოდ ტექნიკური, არამედ ეთიკური მნიშვნელობაც. მკვლევარი პასუხისმგებელია იმაზე, რომ ხელოვნური ინტელექტს არ მიაწოდოს კონფიდენციალური მონაცემები, პერსონალური ინფორმაცია ან ისეთი მასალა, რომლის გავრცელებაც ეწინააღმდეგება კვლევის ეთიკურ სტანდარტებს. ასევე მნიშვნელოვანია, რომ „პრომპტები“ არ იყოს მიმართული მონაცემების მანიპულირების, ცრუ ინფორმაციის გენერირებისა ან ეთიკური წესების გვერდის ავლისკენ.⁴⁶

ამრიგად, თანამედროვე სამეცნიერო სივრცეში „პრომპტების“ ფორმულირება უკვე განიხილება, როგორც კვლევითი კომპეტენციის ნაწილი. ხელოვნური ინტელექტის ეთიკური გამოყენება მოითხოვს არა მხოლოდ ტექნოლოგიური ინსტრუმენტების ცოდნას, არამედ პასუხისმგებლიან, გამჭვირვალე და კეთილსინდისიერ მიდგომას, რათა დაცული იყოს სამეცნიერო ნაშრომის ხარისხი, სანდოობა და აკადემიური ეთიკის პრინციპები.

განსაკუთრებულ მნიშვნელობას იძენს ასევე ხელოვნური ინტელექტის მიერ გენერირებული ტექსტის სამეცნიერო ნაშრომში დამოწმების საკითხი. აღნიშნული მიმართულებით ჩამოყალიბებულია რამდენიმე ძირითადი მიდგომა, რომელსაც მსოფლიოს წამყვანი გამომცემლები და აკადემიური ორგანიზაციები იზიარებენ:

„პუბლიკაციების ეთიკის კომიტეტი“ (COPE) ხაზს უსვამს იმას, რომ ხელოვნური ინტელექტი ვერ იქნება ნაშრომის ავტორი, რადგან მას არ გააჩნია სამართლებრივი და აკადემიური პასუხისმგებლობის ადების უნარი. ამასთან, COPE-ი მოითხოვს, რომ ავტორებმა მკაფიოდ აღწერონ ხელოვნური ინტელექტის გამოყენების ფორმა, ინსტრუმენტი და დანიშნულება ნაშრომის შესაბამის ნაწილში, მაგალითად, „მეთოდოლოგიისა“ და „სამადლობელო“ გვერდებზე.

“Springer Nature”-ის პოლიტიკის მიხედვით, ხელოვნური ინტელექტის გამოყენება დასაშვებია ტექსტის რედაქტირების, გრამატიკული ნორმების შემოწმებისა და გარკვეული კვლევითი პროცესების მხარდასაჭერად, თუმცა იგი ვერ იქნება მითითებული ავტორად. “Springer Nature”-ის რეკომენდაციით მკვლევარებმა დეტალურად უნდა აღწერონ ნაშრომის მეთოდოლოგიურ ნაწილში ხელოვნური ინტელექტის გამოყენების შემთხვევა. ასევე, გამომცემლობა განსაკუთრებულ ყურადღებას უთმობს ხელოვნური ინტელექტის მიერ გენერირებული სურათებისა და ვიზუალური მასალის გამოყენების შეზღუდვას საავტორო უფლებებისა და სანდოობის საკითხების გამო.

“Elsevier”-ის პოლიტიკა ეფუძნება „გამჭვირვალობის უპირატესობის“ მიდგომას. “Elsevier”-ი მოითხოვს, რომ მკვლევარებმა მიუთითონ, თუ რა ტიპის ინსტრუმენტი გამოიყენეს, რა მიზნით და

⁴⁶ Purificato, E., D. Bili, R. Jungnickel, V. Ruiz Serra, J. Fabiani, K. Abendroth Dias, D. Fernandez Llorca, and E. Gomez. 2025. *The Role of Artificial Intelligence in Scientific Research: A Science for Policy, European Perspective*.

ტექსტის რომელ ნაწილზე მოახდინა აღნიშნულმა გავლენა. ამავე დროს, გამომცემლობა მკაფიოდ აცხადებს, რომ ხელოვნური ინტელექტი ვერ ჩაითვლება ავტორად, რადგან ავტორობას თან ახლავს პასუხისმგებლობა ტექსტის სიზუსტესა და სამეცნიერო კეთილსინდისიერებაზე.

“Taylor & Francis” შედარებით უფრო მოქნილ პოლიტიკას ატარებს და მიიჩნევს, რომ ხელოვნური ინტელექტი შეიძლება გამოყენებულ იქნეს იდეების გენერირების, ტექსტის სტრუქტურირების ან ენობრივი დახვეწის მიზნით, თუმცა აუცილებელია მისი გამოყენების გამჟღავნება და საბოლოო ტექსტის სიზუსტეზე სრული პასუხისმგებლობის აღება ავტორის მიერ.

„ამერიკის ფსიქოლოგთა ასოციაცია“ (APA)⁴⁷ ასევე ადგენს, რომ ხელოვნური ინტელექტის გამოყენება სამეცნიერო ნაშრომის მომზადების პროცესში უნდა იყოს მითითებული და შესაბამისი წესით ციტირებული. APA-ს რეკომენდაციებით, ხელოვნური ინტელექტის გამოყენება ძირითადად უნდა აღინიშნოს „მეთოდოლოგიის“ სტრუქტურულ ერთეულში.

„გამოთვლითი მანქანების ასოციაცია“ (ACM)⁴⁸ განმარტავს, რომ ხელოვნური ინტელექტი არ შეიძლება იყოს ავტორი, თუმცა მისი გამოყენება ნებადართულია იმ პირობით, რომ ავტორი სრულად გაამჟღავნებს გამოყენების ფაქტს. ACM-ი განსაკუთრებულ ყურადღებას აქცევს კვლევის გამჭვირვალობისა და პასუხისმგებლობის საკითხს.

შედეგი

სამეცნიერო ნაშრომების შექმნის პროცესში ხელოვნური ინტელექტის გამოყენებამ მნიშვნელოვნად შეცვალა ინფორმაციის მოძიების, ანალიზისა და ტექსტის ფორმირების პრაქტიკა. ამასთან ერთად, განსაკუთრებული მნიშვნელობა შეიძინა აკადემიური კეთილსინდისიერების დაცვის საკითხმა, რადგან ხელოვნური ინტელექტის სისტემების გამოყენება დაკავშირებულია როგორც შესაძლებლობებთან, ისე ეთიკურ რისკებთან. სწორედ ამიტომ, თანამედროვე აკადემიურ სივრცეში ყურადღება ეთმობა არა მხოლოდ ხელოვნური ინტელექტის გამოყენების ფაქტს, არამედ იმასაც, თუ როგორ ხდება მისი გამოყენება კვლევისა და პუბლიკაციის პროცესში.

კვლევის შედეგად დადგინდა, რომ ხელოვნური ინტელექტის გამოყენება სამეცნიერო ნაშრომების ფორმირების პროცესში ერთდროულად ქმნის როგორც მნიშვნელოვან შესაძლებლობებს, ისე სერიოზულ ეთიკურ გამოწვევებს. ხელოვნური ინტელექტის ინსტრუმენტები ეფექტურად გამოიყენება მონაცემთა ანალიზის, ლიტერატურის მოძიების, ტექსტის რედაქტირების, იდეების გენერირებისა და კვლევითი პროცესების ოპტიმიზაციისთვის, რაც მნიშვნელოვნად ამცირებს დროით და ტექნიკურ რესურსებს. მიუხედავად ამისა, კვლევამ აჩვენა, რომ ხელოვნური ინტელექტის გამოყენება პირდაპირ არის დაკავშირებული აკადემიური კეთილსინდისიერების დაცვის აუცილებლობასთან.

გამოიკვეთა, რომ ხელოვნური ინტელექტის სისტემები არ არის სრულად დაცული შემთხვევითი და სისტემური შეცდომებისგან. განსაკუთრებით პრობლემურია ყალბი ციტატების, არაზუსტი

⁴⁷ American Psychological Association.

⁴⁸ Association for Computing Machinery.

ინფორმაციის, მიკერძოებული მონაცემებისა და პლაგიატის რისკი. აღნიშნული გარემოება საფრთხეს უქმნის სამეცნიერო ნაშრომის სანდოობასა და ვალიდურობას, თუ მკვლევარი სათანადოდ არ გადაამოწმებს ხელოვნური ინტელექტის მიერ გენერირებულ ინფორმაციას. შესაბამისად, კვლევამ დაადასტურა, რომ საბოლოო პასუხისმგებლობა სამეცნიერო ნაშრომის სიზუსტესა და ეთიკურ შესაბამისობაზე ყოველთვის ეკისრება მკვლევარს.

კვლევის ფარგლებში ასევე დადგინდა, რომ საერთაშორისო აკადემიური ორგანიზაციებისა და გამოცემლობების მიდგომები თანხმდება რამდენიმე ძირითად პრინციპზე, როგორებიცაა: გამჭვირვალობა, ანგარიშვალდებულება, პასუხისმგებლობა და ავტორობის ადამიანური ბუნება. “COPE”-ის, “Springer Nature”-ის, “Elsevier”-ის, “Taylor & Francis”-ის, APA-სა და ACM-ის პოლიტიკების ანალიზმა აჩვენა, რომ ხელოვნური ინტელექტი არ შეიძლება ჩაითვალოს სამეცნიერო ნაშრომის ავტორად, თუმცა მისი გამოყენების ფაქტი მკაფიოდ და გამჭვირვალედ უნდა იყოს მითითებული ნაშრომში.

კვლევის ერთ-ერთ მნიშვნელოვან მიგნებად გამოიკვეთა „პრომპტების“ სწორად ფორმულირების მნიშვნელობა. დადგინდა, რომ ხელოვნური ინტელექტის მიერ გენერირებული ინფორმაციის ხარისხი, სიზუსტე და ეთიკური შესაბამისობა მნიშვნელოვნად არის დამოკიდებული მკვლევრის მიერ დასმული მოთხოვნის სიზუსტეზე. ბუნდოვანი ან არაზუსტად ფორმულირებული „პრომპტები“ ზრდის არასწორი, არასრული ან ეთიკურად პრობლემური ინფორმაციის მიღების რისკს.

კვლევამ ასევე აჩვენა, რომ ეთიკური რისკების შემცირებისთვის აუცილებელია როგორც სამართლებრივი და ინსტიტუციური რეგულაციების განვითარება, ისე ტექნოლოგიური და საგანმანათლებლო მექანიზმების დანერგვა. განსაკუთრებული მნიშვნელობა ენიჭება მკვლევართა ცნობიერების ამაღლებას, ხელოვნური ინტელექტის ეთიკის სწავლების გაძლიერებას, უნივერსიტეტებისა და სამეცნიერო ჟურნალების მიერ მკაფიო სახელმძღვანელო პრინციპების შემუშავებას.

დასკვნა

ამრიგად, ხელოვნური ინტელექტი თანამედროვე სამეცნიერო კვლევის მნიშვნელოვანი ინსტრუმენტია, რომელმაც არსებითად შეცვალა კვლევის, ანალიზისა და აკადემიური ტექსტების ფორმირების პროცესი. მიუხედავად მისი მრავალმხრივი შესაძლებლობებისა, კვლევამ აჩვენა, რომ ხელოვნური ინტელექტის გამოყენება სამეცნიერო ნაშრომებში აუცილებლად უნდა ეფუძნებოდეს აკადემიური კეთილსინდისიერების, გამჭვირვალობისა და პასუხისმგებლობის პრინციპებს. ხელოვნური ინტელექტი ვერ ჩაანაცვლებს მკვლევრის ინტელექტუალურ, ეთიკურ და პროფესიულ პასუხისმგებლობას, რადგან მას არ გააჩნია მორალური განსჯის, ანგარიშვალდებულებისა და სამართლებრივი პასუხისმგებლობის უნარი.

კვლევის საფუძველზე შეიძლება დასკვნის სახით ითქვას, რომ სამეცნიერო სივრცეში აუცილებელია ხელოვნური ინტელექტის გამოყენების ერთიანი ეთიკური სტანდარტებისა და პრაქტიკული სახელმძღვანელოების განვითარება, რაც უზრუნველყოფს აკადემიური ხარისხისა და სამეცნიერო ნდობის შენარჩუნებას. ხელოვნური ინტელექტის გამოყენება უნდა

განიხილებოდეს, როგორც დამხმარე ინსტრუმენტი და არა, როგორც ავტონომიური შემოქმედებითი ან სამეცნიერო სუბიექტი.

კვლევის საფუძველზე მიზანშეწონილია შემდეგი რეკომენდაციების შემუშავება:

- სამეცნიერო ჟურნალებმა, უნივერსიტეტებმა და კვლევითმა ინსტიტუტებმა უნდა შეიმუშაონ ხელოვნური ინტელექტის გამოყენების მკაფიო, ერთიანი და აღსრულებადი ეთიკური პოლიტიკა.
- მკვლევრებმა სავალდებულოდ უნდა მიიჩნიონ ხელოვნური ინტელექტის გამოყენების გამჭვირვალედ მითითება ნაშრომის შესაბამის ნაწილში, განსაკუთრებით მაშინ, როდესაც იგი გამოიყენება ტექსტის გენერირების, რედაქტირების ან მონაცემთა ანალიზის მიზნით.
- აუცილებელია ხელოვნური ინტელექტის მიერ გენერირებული ინფორმაციის სიზუსტის, ციტატებისა და წყაროების სავალდებულო გადამოწმება მკვლევრის მიერ.
- უნივერსიტეტებსა და სამეცნიერო პროგრამებში უნდა გაძლიერდეს ხელოვნური ინტელექტის ეთიკის სწავლება, რათა მკვლევარებმა შეძლონ ტექნოლოგიების პასუხისმგებლიანად გამოყენება.
- განსაკუთრებული ყურადღება უნდა დაეთმოს „პრომპტების“ სწორ ფორმულირებას, რადგან სწორედ იგი განსაზღვრავს მიღებული ინფორმაციის ხარისხსა და ეთიკურ შესაბამისობას.
- აუცილებელია კონფიდენციალური და პერსონალური მონაცემების დაცვის მაღალი სტანდარტების უზრუნველყოფა ხელოვნური ინტელექტის სისტემების გამოყენების პროცესში.
- სამეცნიერო საზოგადოებამ უნდა გააძლიეროს ტექნოლოგიური მექანიზმების გამოყენება, რომლებიც შესაძლებელს გახდის ხელოვნური ინტელექტის მიერ გენერირებული ტექსტის იდენტიფიცირებასა და კონტროლს.
- ხელოვნური ინტელექტის გამოყენება უნდა ემსახურებოდეს მკვლევრის ანალიტიკური და კვლევითი შესაძლებლობების გაძლიერებას და არა სამეცნიერო ნაშრომების ფორმალურ გენერირებას.

ყოველივე ზემოხსენებულიდან გამომდინარე, ხელოვნური ინტელექტის ეთიკური და პასუხისმგებლიანი გამოყენება წარმოადგენს თანამედროვე აკადემიური სივრცის ერთ-ერთ უმნიშვნელოვანეს გამოწვევასა და ამავდროულად განვითარების შესაძლებლობას. მხოლოდ კეთილსინდისიერების პრინციპის, გამჭვირვალობისა და პროფესიული პასუხისმგებლობის დაცვით იქნება შესაძლებელი ხელოვნური ინტელექტის ინტეგრირება სამეცნიერო კვლევაში ისე, რომ კვლევის პროცესში შენარჩუნდეს მეცნიერების სანდოობა, ხარისხი და აკადემიური ღირებულებები.

ბიბლიოგრაფია

უცხოური წყაროები

1. Alvarado, R., 2022, *Computer Simulations as Scientific Instruments*, Foundations of Science, 27(3), p. 1183–1205.

2. Alvarado, R., 2023, *AI as an Epistemic Technology*, Science and Engineering Ethics, 29, 32.
3. Ball, P., 2023, *Is AI Leading to a Reproducibility Crisis in Science?*, Nature, 624, p. 22–25.
4. Bhattacharyya, M., Miller, V.M., Bhattacharyya, D., Miller, L.E., 2023, *High Rates of Fabricated and Inaccurate References in ChatGPT-Generated Medical Content*, Cureus, 15(5), e39238.
5. Caron, M.M., Dohan, S.B., Barnes, M., Bierer, B.E., 2023, *Defining “Recklessness” in Research Misconduct Proceedings*, Accountability in Research, p. 1–23.
6. Daneshjou, R., Smith, M.P., Sun, M.D., Rotemberg, V., Zou, J., 2021, *Lack of Transparency and Potential Bias in Artificial Intelligence Data Sets and Algorithms: A Scoping Review*, JAMA Dermatology, 157(11), p. 1362–1369.
7. Deanda, D., Alsmadi, I., Guerrero, J., Liang, G., 2025, *Defending Mutation-Based Adversarial Text Perturbation: A Black-Box Approach*, Cluster Computing, 28, 196.
8. Floridi, L., Cowls, J., 2022, *A Unified Framework of Five Principles for AI in Society*, in book: Silvio Carta (ed.), *Machine Learning and the City: Applications in Architecture and Urban Design*, John Wiley & Sons.
9. Gandhi, S., Ahmed, M.S., Huang, L., Behl, A., Manrai, A.K., 2025, *AI’s Impact on Academic Research Integrity: Guidelines from the Editors of Journal of Global Marketing*, Journal of Global Marketing, 38(3), p. 189–192.
10. Garin, S.P., Parekh, V.S., Sulam, J., Yi, P.H., 2023, *Medical Imaging Data Science Competitions Should Report Dataset Demographics and Evaluate for Bias*, Nature Medicine, 29(5), p. 1038–1039.
11. Giere, R., Bickle, J., Maudlin, R.F., 2005, *Understanding Scientific Reasoning*, 5th edition, Belmont, Wadsworth.
12. Gunning, D., Aha, D., 2019, *DARPA’s Explainable Artificial Intelligence (XAI) Program*, AI Magazine, 40(2), p. 44–58.
13. Haack, S., 2007, *Defending Science within Reason*, New York, Prometheus Books.
14. Howson, C., Urbach, P., 2005, *Scientific Reasoning: A Bayesian Approach*, 3rd edition, New York, Open Court.
15. Kitcher, P., 1993, *The Advancement of Knowledge*, New York, Oxford University Press.
16. Licht, K., Licht, J., 2020, *Artificial Intelligence, Transparency, and Public Decision-Making: Why Explanations Are Key When Trying to Produce Perceived Legitimacy*, AI & Society, 35, p. 917–926.
17. Llerena-Izquierdo, J., Ayala-Carabajo, R., 2025, *Ethics of the Use of Artificial Intelligence in Academia and Research: The Most Relevant Approaches, Challenges and Topics*, Informatics, 12(4), 111.
18. Merton, R., 1973, *The Sociology of Science*, Chicago, University of Chicago Press.
19. Mitchell, M., 2019, *Artificial Intelligence*, New York, Picador.
20. Office of Science and Technology Policy, 2000, *Federal Research Misconduct Policy*, Federal Register, 65(235), p. 76260–76264.
21. Purificato, E., Bili, D., Jungnickel, R., Ruiz Serra, V., Fabiani, J., Abendroth Dias, K., Fernandez Llorca, D., Gomez, E., 2025, *The Role of Artificial Intelligence in Scientific Research: A Science for Policy European Perspective*.

22. Resnik, D.B., Elliott, K.C., 2019, *Value-Entanglement and the Integrity of Scientific Research*, Studies in History and Philosophy of Science, 75, p. 1–11.
23. Resnik, D.B., Hosseini, M., 2025, *The Ethics of Using Artificial Intelligence in Scientific Research: New Guidance Needed for a New Tool*, AI Ethics, 5(2), p. 1499–1521.
24. Shamoo, A.E., Resnik, D.B., 2022, *Responsible Conduct of Research*, 4th edition, New York, Oxford University Press.
25. Shapin, S., 1995, *Here and Everywhere: Sociology of Scientific Knowledge*, Annual Review of Sociology, 21, p. 289–321.
26. Walters, W.H., Wilder, E.I., 2023, *Fabrication and Errors in the Bibliographic Citations Generated by ChatGPT*, Scientific Reports, 13, 14045.

დამატებითი წყაროები

1. International Committee of Medical Journal Editors, 2023, *Recommendations for the Conduct, Reporting, Editing, and Publication of Scholarly Work in Medical Journals*, ხელმისაწვდომია: <https://www.icmje.org/icmje-recommendations.pdf> (წვდომის თარიღი: 18.05.2026).
2. Milmo, D., 2023, *Two US Lawyers Fined for Submitting Fake Court Citations from ChatGPT*, The Guardian, ხელმისაწვდომია: <https://www.theguardian.com/technology/2023/jun/23/two-us-lawyers-fined-submitting-fake-court-citations-chatgpt> (წვდომის თარიღი: 18.05.2026).
3. National Institutes of Health, 2023, *Guidelines for the Conduct of Research in the Intramural Program of the NIH*, ხელმისაწვდომია: https://oir.nih.gov/system/files/media/file/2023-11/guidelines-conduct_research.pdf (წვდომის თარიღი: 18.05.2026).
4. The European Code of Conduct for Research Integrity, Revised Edition 2023, 2023, ხელმისაწვდომია: <https://allea.org/code-of-conduct/> (წვდომის თარიღი: 18.05.2026).
5. UNESCO, 2024, *Ethics of Artificial Intelligence*, ხელმისაწვდომია: <https://www.unesco.org/en/artificial-intelligence/recommendation-ethics> (წვდომის თარიღი: 18.05.2026).
6. Gandhi, S., Ahmed, M.S., Huang, L., Behl, A., Manrai, A.K., 2025, *AI's Impact on Academic Research Integrity: Guidelines from the Editors of Journal of Global Marketing*, Journal of Global Marketing, 38(3), p. 189–192, ხელმისაწვდომია: <https://doi.org/10.1080/08911762.2025.2488060> (წვდომის თარიღი: 18.05.2026).
7. Gunning, D., Aha, D., 2019, *DARPA's Explainable Artificial Intelligence (XAI) Program*, AI Magazine, 40(2), p. 44–58, ხელმისაწვდომია: <https://doi.org/10.1609/aimag.v40i2.2850> (წვდომის თარიღი: 18.05.2026).
8. Llerena-Izquierdo, J., Ayala-Carabajo, R., 2025, *Ethics of the Use of Artificial Intelligence in Academia and Research: The Most Relevant Approaches, Challenges and Topics*, Informatics, 12(4), 111, ხელმისაწვდომია: <https://doi.org/10.3390/informatics12040111> (წვდომის თარიღი: 18.05.2026).
9. Alvarado, R., 2023, *AI as an Epistemic Technology*, Science and Engineering Ethics, 29, 32, ხელმისაწვდომია: <https://link.springer.com/article/10.1007/s11948-023-00451-3> (წვდომის თარიღი: 18.05.2026).