

UDC 2213

SCOPUS CODE 2213

<https://doi.org/10.36073/1512-0996-2025-3-309-316>

შრომის უსაფრთხოების ახალი ეტაპი – ხელოვნური ინტელექტის ინტეგრაცია სამუშაო გარემოში

**ნიკოლოზ
შავერდაშვილი**

შრომის უსაფრთხოებისა და საგანგებო მდგომარეობის მართვის დეპარტამენტი,
საქართველოს ტექნიკური უნივერსიტეტი, საქართველო, 0160, თბილისი,
მ. კოსტავას 75
E-mail: nikoloz298@gmail.com

ნანა რაზმაძე

შრომის უსაფრთხოებისა და საგანგებო მდგომარეობის მართვის დეპარტამენტი,
საქართველოს ტექნიკური უნივერსიტეტი, საქართველო, 0160, თბილისი,
მ. კოსტავას 75
E-mail: n.razmadze@gtu.ge

ნინო რატიანი

შრომის უსაფრთხოებისა და საგანგებო მდგომარეობის მართვის დეპარტამენტი,
საქართველოს ტექნიკური უნივერსიტეტი, საქართველო, 0160, თბილისი,
მ. კოსტავას 75
E-mail: n.ratiani@gtu.ge

რეცენზენტები:

თ. კუნჭულია, სტუ-ის სამთო-გეოლოგიური ფაკულტეტის პროფესორი

E-mail: t.kunchulia@gtu.ge

დ. ბუცხრიკიძე, სტუ-ის სატრანსპორტო სისტემების და მექანიკის ინჟინერიის ფაკულტეტის ასოცირებული პროფესორი

E-mail: d.butskhrikidze@gtu.ge

ანოტაცია. კვლევა აჩვენებს, რომ ხელოვნური ინტელექტი არა მხოლოდ ტექნოლოგიური ინოვაცია, არამედ მენეჯმენტის სტრატეგიული ინსტრუმენტიცაა. მოცემულ სტატიაში განიხილება ხელოვნური ინტელექტის (AI) როლი კომპანიის თანამშრომელთა შრომის უსაფრთხოების გაუმჯობესების მიმართულებით. წარმოდგენილია თანამედროვე ტექნოლოგიური მიდგომები, რისკების პროგ-

ნოზირების მეთოდები, რეალურ დროში მონიტორინგის სისტემები და ციფრული ტრენინგების მოდელები. მოყვანილია წამყვანი საერთაშორისო კომპანიების მაგალითები.

საკვანძო სიტყვები: მონაცემთა ანალიზი; რისკების მენეჯმენტი; ტექნოლოგიური პროგრესი; შრომის უსაფრთხოება; ხელოვნური ინტელექტი.

შესავალი

ხელოვნური ინტელექტი (ინგლ. Artificial Intelligence, შემოკლებით **AI**) არის ტექნოლოგია, რომელიც საშუალებას აძლევს კომპიუტერებსა და სხვა მოწყობილობებს შეასრულონ ისეთი ამოცანები, რომელიც მიზნად ისახავს ადამიანის ინტელექტის გაგებას და ჩვეულებრივ, მოითხოვს ადამიანის ინტელექტს. მაგალითად:

- ინფორმაციის გაანალიზება და დასკვნების გამოტანა;
- საუბრის ან ტექსტის გაგება და გენერირება;
- გამოსახულებების ამოცნობა;
- გადაწყვეტილებების მიღება;
- სწავლა წარსული გამოცდილებიდან (ე.წ. „machine learning“).

ხელოვნური ინტელექტის (AI) იდეა ჯერ კიდევ მე-20 საუკუნის დასაწყისში გაჩნდა, მაგრამ AI-ის ოფიციალური დაბადების წლად 1956 წელი ითვლება.

1956 წელი – ტერმინის წარმოშობა:

- აშშ-ში, დარტმუთის უნივერსიტეტში ჩატარდა სამეცნიერო სამუშაო სემინარი, სადაც პირველად გაჩნდა ტერმინი “ხელოვნური ინტელექტი” (Artificial Intelligence);
- სემინარის ორგანიზატორი იყო ჯონ მაკკარტი (John McCarthy) — მას ხშირად “ხელოვნური ინტელექტის მამად” მიიჩნევენ.

XX ს-ის 50–60-იანები:

- შეიქმნა პირველი პრიმიტიული AI პროგრამები (მაგალითად, ლოგიკური ამოცანების გადაჭრისთვის);

- ალან ტიურინგმა შემოგვთავაზა “ტიურინგის ტესტი”, რომლითაც შეიძლება შეფასდეს, რამდენად “ინტელექტუალურია” მანქანა.

1997 წელი:

- IBM-ის სუპერკომპიუტერმა Deep Blue-მ სძლია მსოფლიოს ჩემპიონს ჭადრაკში – გარი კასპაროვს.

XXI ს-ის 10-იანები – თანამედროვე AI-ის აღზევება:

- დიდი მონაცემების (Big Data), ძლიერი კომპიუტერების და ღრმა სწავლის (Deep Learning) წყალობით AI-მ მკვეთრი წინსვლა განიცადა;
- Google, Facebook, Microsoft და სხვა გიგანტები აქტიურად ჩაერთვნენ AI-ის განვითარებაში.

AI-ის ფუძემდებელი ეტაპი 1956 წელს დაიწყო, თუმცა მისი სერიოზული გამოყენება და პოპულარობა ბოლო 10–15 წელს განახლებულმა ტექნოლოგიებმა რეალურად შესაძლებელი გახადა.

ძირითადი ნაწილი

ხელოვნური ინტელექტი (AI - Artificial Intelligence) XXI საუკუნის ერთ-ერთ ყველაზე გარდამტეხ ტექნოლოგიურ მიღწევად იქცა. მისი არსი არის ის, რომ სისტემას შეუძლია ადამიანის ინტელექტუალური ფუნქციების – ანალიზის, სწავლის, გადაწყვეტილების მიღებისა და პრობლემის გადაჭრის იმიტაცია. მიუხედავად იმისა, რომ AI-ის ძირითადი კონცეფცია რამდენიმე ათწლეულს ითვლის, მისი პრაქტიკული გამოყენება და გავლენა ეკონომიკაზე, სოციალურ სტრუქტურებსა და ადამიანის ჯანმრთელობაზე განსაკუთრებით თვალსაჩინო გახდა ბოლო წლების განმავლობაში.

თანამედროვე საწარმოო და მომსახურების სექტორებში ადამიანური რესურსის უსაფრთხოება და ჯანმრთელობის დაცვა ერთ-ერთ ცენტრალურ გამოწვევად რჩება. ტექნოლოგიური პროგრესი, განსაკუთრებით ხელოვნური ინტელექტის (AI) ინტეგრაცია, ცვლის უსაფრთხოების და ჯანმრთელობის მართვის ტრადიციულ მოდელებს. AI-ის საშუალებით შესაძლებელი გახდა რეალურ დროში მონაცემთა ანალიზი, რისკების პროგნოზირება და დროული რეაგირება, რაც მნიშვნელოვნად ამცირებს უზედური შემთხვევების, გადაღლილობის და სტრესული ფაქტორების რისკებს. იმ პირობებში, როცა სამუშაო გარემოში ადამიანთა ჯანმრთელობაზე მოქმედი ფაქტორები სულ უფრო კომპლექსური ხდება, ხელოვნური ინტელექტი იქცა ეფექტური მართვისა და პრევენციის მძლავრ ინსტრუმენტად.

გარდა ამისა, შრომის უსაფრთხოების საკითხები უკვე აღარ შემოიფარგლება მხოლოდ ფიზიკური საფრთხეებით. მზარდია ფსიქიკური ჯანმრთელობის, ბიოეთიკისა და პერსონალური მონაცემების დაცვის რისკებიც, რაც კიდევ უფრო ზრდის AI-სადმი ინტერესს როგორც დამსაქმებელთა, ისე სახელმწიფო ინსტიტუტების მხრიდან. შესაბამისად, ამ თემაზე კვლევა არა მხოლოდ პრაქტიკულად, არამედ პოლიტიკურად და სამართლებრივად აქტუალურია.

ხელოვნური ინტელექტის ერთ-ერთი ძლიერი მხარეა მისი მრავალდარგობრივი ინტეგრაციის უნარი. ის გამოიყენება როგორც სამრეწველო პროცესებში, ისე სამართლებრივ ანალიზში, განათლებაში და ყველაზე მნიშვნელოვანი – ადამიანზე ორიენტირებულ დარგებში, როგორიცაა შრომის უსა-

ფრთხოება და ჯანმრთელობის დაცვა. ტექნოლოგიური ტრანსფორმაციის პროცესში AI ხორციელდება როგორც ცალკეული ავტომატიზებული ელემენტების სახით (მაგალითად, რობოტები საწარმოო ხაზზე), ისე როგორც კომპლექსური ანალიტიკური მოდელები, რომლებიც ახდენენ რისკების პროგნოზირებასა და ადრეულ გაფრთხილებას [1].

ამ სისტემების მოქმედება ეფუძნება მონაცემებზე დაყრდნობილ ფუნქციონირებას – რაც უფრო მაღალია ინფორმაციის ხარისხი და მოცულობა, მით უფრო ზუსტია პროგნოზი ან გადაწყვეტილება. ეს მახასიათებელი განსაკუთრებით მნიშვნელოვანია შრომის უსაფრთხოების სფეროში, სადაც განუწყვეტელი მონიტორინგი და დროული რეაგირება განსაზღვრავს სიცოცხლისა და ჯანმრთელობის შენარჩუნების შესაძლებლობას. მაგალითად, მანქანური სწავლის მოდელები შეძლებს საფრთხის სწრაფად იდენტიფიცირებას და რეაგირებას ისე, რომ მინიმუმამდე დაიყვანოს ადამიანური ფაქტორიდან გამომდინარე შეცდომები.

შრომის უსაფრთხოება ყოველთვის წარმოადგენდა ერთ-ერთ მნიშვნელოვან გამოწვევას მრეწველობისა და მძიმე შრომის სეგმენტებში, სადაც ადამიანის ჯანმრთელობა და სიცოცხლე პირდაპირ არის დამოკიდებული გარემოს რისკებზე. ამ გამოწვევის საპასუხოდ, თანამედროვე ტექნოლოგიური პროგრესი და ხელოვნური ინტელექტის (AI) ინტეგრაცია არსებითად ცვლის უსაფრთხოების მექანიზმების მიდგომებს. იმის ნაცვლად, რომ უსაფრთხოების უზრუნველყოფა ემყარებოდა მხოლოდ ადამიანის დაკვირვებასა და გამოცდილებას, დღეს რეალურ დროში მოქმედი AI სისტემები ქმნიან გარემოს, სადაც საფრთხე დროულად აღმოჩენი-

ლია, შეფასებულია და შესაძლებელია მისი პრევენცია.

AI-ის მიერ შრომის უსაფრთხოების მხარდაჭერა იწყება მონაცემების შეგროვებით. თანამედროვე სენსორები და ვიდეოანალიტიკური სისტემები, რომლებსაც მართავს მანქანური სწავლება, აკვირდება სამუშაო გარემოს, აღიქვამს მოძრაობის ტრაექტორიებს, აიდენტიფიცირებს არასწორ მოქმედებებს და ააქტიურებს გაფრთხილებას ჯერ კიდევ მანამ, სანამ მოხდება რეალური დაზიანება. ეს მექანიზმები მნიშვნელოვნად ამცირებს რეაქციის დროს და უზრუნველყოფს იმ წერტილოვან რეაგირებას, რასაც ადამიანის კონტროლი ყოველთვის ვერ ახერხებს. ამ კონტექსტში განსაკუთრებით მნიშვნელოვანია კავშირი ავტომატიზაციასა და ადამიანურ კონტროლს შორის – სისტემა არ ცვლის ადამიანს, არამედ ეხმარება მას, რათა მეტი სიზუსტით იმოქმედოს.

მნიშვნელოვანი უპირატესობა AI-ის გამოყენებისას არის რისკების პროგნოზირება და სტატისტიკური ანალიზი. დიდი რაოდენობით ისტორიული მონაცემების დამუშავებით, AI ალგორითმები აფასებენ, სად და როდის არის უბედური შემთხვევის ყველაზე დიდი ალბათობა. ეს შესაძლებელს ხდის პრევენციული ქმედებების დაგეგმვას, მაგალითად, ისეთი ამოცანების გადანაწილებას, რომლებიც დამქირავებლისთვის ცნობილია როგორც მაღალი რისკის შემცველი. ეს მიდგომა ემყარება არა მხოლოდ ტექნოლოგიურ ინოვაციას, არამედ მტკიცებულებებზე დაფუძნებულ მენეჯმენტს [1].

AI ასევე იძლევა შესაძლებლობას საწარმოებში შეიქმნას ციფრული ტრენინგები და სიმულაციები,

სადაც თანამშრომლები სწავლობენ რეალურ გარემოში მოქმედების წესებს, რაც ამცირებს მომავალი შეცდომების ალბათობას და ზრდის უსაფრთხოების კულტურას ორგანიზაციაში. ეს ტრენინგები ხშირად გამოიყენება მძიმე მრეწველობის, ენერგეტიკის და ლოგისტიკის სფეროებში, სადაც მცირე შეცდომასაც კი შესაძლოა მოჰყვეს სერიოზული შედეგი [13].

ხელოვნური ინტელექტის ინტეგრაცია შრომის უსაფრთხოებასა და ჯანმრთელობის დაცვაში არ არის მხოლოდ თეორიული მოდელი — ის უკვე რეალურად ფუნქციონირებს მსოფლიოს წამყვან კომპანიებში.

Amazon-ი იყენებს AI ტექნოლოგიებს დასაქმებულთა ფიზიკური დატვირთვის მონიტორინგისთვის. სპეციალური სენსორები და მონაცემთა ანალიზის მოდელები აფასებენ თითოეული თანამშრომლის ყოველდღიურ მოძრაობასა და ძალისხმევას, რათა გამოირიცხოს გადაჭარბებული დატვირთვა, რაც ხშირად იწვევს ქრონიკულ ტკივილებსა და დაზიანებებს [1]. ეს ტექნოლოგია კომპანიისთვის დისკომფორტისა და უმოქმედობის შემცირებას, ხოლო დასაქმებულისთვის ჯანმრთელობის შენარჩუნებას ნიშნავს.

Soter Analytics ხელოვნურ ინტელექტზე დამყარებული ვიდეო-აუდიო სამეტვალყურეო მოწყობილობის კომპანიის მეშვეობით (პერსონალური მოწყობილობები + გამოსახულების ანალიზი) თანამშრომლები აერთიანებენ OSHA-ს წესების შესრულებას; 90-ზე მეტი კომპანია იყენებს მათ გადაწყვეტას (IKEA, Coca-Cola, DHL).

ჯანდაცვის სექტორშიც ვხვდებით წარმატებულ

მაგალითებს. კლინიკებში, როგორცაა **Mayo Clinic** და **Cleveland Clinic**, AI გამოიყენება როგორც კლინიკური გადაწყვეტილების მხარდამჭერი სისტემა. თანამშრომლების ჯანმრთელობის მონაცემებზე დაყრდნობით, სისტემა აფასებს რისკებს – მაგალითად, გულის შეტევის ან დიაბეტის განვითარების ალბათობას და თანამშრომლებს სთავაზობს წინასწარ სამედიცინო კონსულტაციებს, ან საჭირო სამოქმედო გეგმის ფორმირებას [11].

ტექნოლოგიური პროგრესი არ არის თავისთავად მიზანი, – ის არის მექანიზმი, რომელიც უნდა ემსახურებოდეს ადამიანის კეთილდღეობას, ღირსებას და უსაფრთხოებას.

შეიძლება ითქვას, რომ ხელოვნური ინტელექტი ქმნის მრავალმხრივ შესაძლებლობებს როგორც შრომის უსაფრთხოების, ისე ადამიანის ჯანმრთელობის დაცვის სფეროში. რეალურ დროში მონაცემთა შეგროვება, სიღრმისეული ანალიზი, ავტომატური გაფრთხილებები და რისკების პროგნოზირება განაპირობებს იმას, რომ AI უკვე იქცა გადამწყვეტ ფაქტორად მუშაკთა სიცოცხლისა და ჯანმრთელობის შენარჩუნებაში.

ტექნოლოგია ხელს უწყობს ტრენინგების გაუმჯობესებას, პროცესების ოპტიმიზაციას და სამუშაო კულტურის ამაღლებას. წარმატებული კორპორატიული მაგალითების ანალიზი ცხადყოფს, რომ AI არ არის მხოლოდ ტექნიკური ინოვაცია, ის წარმოადგენს კომპანიების მენეჯმენტის სტრატეგიულ კომპონენტს.

დასკვნა

განხილულია ხელოვნური ინტელექტის მნიშვნელობა შრომის უსაფრთხოებასა და ადამიანის

ჯანმრთელობის დაცვაში. სხვადასხვა თეორიული, პრაქტიკული და სამართლებრივი კუთხით გაანალიზდა AI-ის ინტეგრაციის რეალური შედეგები, პერსპექტივები და გამოწვევები. კვლევამ აჩვენა, რომ ტექნოლოგიური პროგრესი არ არის თავისთავად მიზანი – ის არის მექანიზმი, რომელიც უნდა ემსახურებოდეს ადამიანის კეთილდღეობას, ღირსებას და უსაფრთხოებას.

შეიძლება ითქვას, რომ ხელოვნური ინტელექტი ქმნის მრავალმხრივ შესაძლებლობას როგორც შრომის უსაფრთხოების, ისე ადამიანის ჯანმრთელობის დაცვის სფეროში. რეალურ დროში მონაცემთა შეგროვება, სიღრმისეული ანალიზი, ავტომატური გაფრთხილებები და რისკების პროგნოზირება განაპირობებს იმას, რომ AI უკვე იქცა გადამწყვეტ ფაქტორად მუშაკთა სიცოცხლისა და ჯანმრთელობის შენარჩუნებაში. ამასთანავე, ტექნოლოგია ხელს უწყობს ტრენინგების გაუმჯობესებას, პროცესების ოპტიმიზაციას და სამუშაო კულტურის ამაღლებას. წარმატებული კორპორატიული მაგალითების ანალიზი ცხადყოფს, რომ AI არ არის მხოლოდ ტექნიკური ინოვაცია, ის წარმოადგენს მენეჯმენტის სტრატეგიულ კომპონენტს. თუმცა, ნაშრომმა ასევე დაადასტურა, რომ ხელოვნური ინტელექტის გამოყენებას თან სდევს სერიოზული ეთიკური, სამართლებრივი და სოციალური კითხვები. განსაკუთრებით კრიტიკულია პერსონალური მონაცემების დაცვა, პასუხისმგებლობის გამიჯვნა ტექნიკური შეცდომების შემთხვევაში და ტექნოლოგიაზე თანაბარი წვდომა სხვადასხვა სეგმენტისთვის. თუ AI სისტემები არ იქნება გამჭვირვალე და სამართლებრივად გამართული, მათი პოტენციალი შესაძლოა გადაიქ-

ცეს არა დაცვის, არამედ კონტროლისა და ჩაგვრის ინსტრუმენტად.

- AI-ის აქვს პოტენციური მნიშვნელოვნად გააუმჯობესოს შრომის უსაფრთხოება და ჯანმრთელობის დაცვა, მაგრამ მხოლოდ მაშინ, როცა მას თან ახლავს მკაფიო სამართლებრივი ჩარჩო და ეთიკური მიდგომა.

- წარმატებული პრაქტიკა აჩვენებს, რომ ტექნოლოგიის ინტეგრაცია უშუალოდ კავშირშია მენეჯმენტის ხედვასთან, თანამშრომელთა ჩართულობასთან და განათლებასთან.
- ტექნოლოგიამ კი არ უნდა ჩაანაცვლოს ადამიანი, არამედ უნდა გააძლიეროს მისი შესაძლებლობები, ეს არის ეფექტური თანამშრომლობის მოდელი.

ლიტერატურა

1. Aldosari, B., Aldosari, H., & Alanazi, A. (2025). Challenges of artificial intelligence in medicine. In *Building the future of health intelligence* (Vol. 308, pp. 230–237). IOS Press.
<https://ebooks.iospress.nl/doi/10.3233/SHTI250039>
2. Ausaf, H., & Khan, M. (2025). Threats by artificial intelligence to human health and human existence. *Human Rights Law Review*. <https://humanrightlawreview.in/wp-content/uploads/2025/03/Threats-by-Artificial-Intelligence-to-Human-Health-and-Human-Existence.pdf>
3. EBU PVSYST. (2025). Risk assessments in occupational health and safety for power substations. *Conference Proceedings*. <https://search.proquest.com/openview/69c346020e7b7fde1c768d29a85be94d/1?pq-origsite=gscholar&cbl=1536338>
4. FARMACEUTICE, Ș. (2025). Legal perspectives on AI and risk in public health. *Științe Juridice și Administrative*. <https://search.proquest.com/openview/18b1551adc487ccee747a2b15719767f/1?pq-origsite=gscholar&cbl=2026879>
5. Gomase, V. S., Ghatule, A. P., & Sharma, R. (2025). Cybersecurity and compliance in clinical trials: The role of artificial intelligence in secure healthcare management. *PubMed*.
<https://pubmed.ncbi.nlm.nih.gov/40277117/>
6. Klein, S., & Campisi, L. M. (2025). From over the pond: The European Union's comprehensive AI legislation comes to America. *The Brief*.
<https://search.proquest.com/openview/feb08e459ba892b8936dd6c2b5e0fed1/1?pq-origsite=gscholar&cbl=17556>
7. Pandas, S., Sasneha, S., & Nayak, P. (2025). Role of enhanced visual cryptography algorithm in cybersecurity. In *Machine learning and artificial intelligence* (pp. 141–158). Springer.
<https://books.google.com/books?hl=en&lr=&id=YY5ZEQAQBAJ&oi=fnd&pg=PA141>
8. Papasani, P. K., & Papasani, R. T. (2025). The use of artificial intelligence and data analytics in nutrition. *ResearchGate*. <https://www.researchgate.net/profile/Pavan-Kumar->

- Papasani/publication/390907621_The_use_of_Artificial_Intelligence_and_data_analytics_in_nutrition/links/6802716ebf974b23ab3f6c/The-use-of-Artificial-Intelligence-and-data-analytics-in-nutrition.pdf
9. Ren, P., Chen, Z., Chen, S., Lu, Y., & Chen, Q. (2025). Integration of Raman spectroscopy and machine learning for highly efficient identification of plastics. *SSRN Electronic Journal*.
https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5194565
 10. Singh, S., Kaur, P., Kaur, I., Singh, G., & Kaur, S. (2025). A predictive framework using advanced machine learning approaches for measuring and analyzing the impact of synthetic agrochemicals on human health. *Scientific Reports*. <https://www.nature.com/articles/s41598-025-00509-1>
 11. Somireddy, S. (2025). AI-driven diagnostics: The role of machine learning in healthcare. *International Medical Research Journal Review*. <https://imrjr.com/wp-content/uploads/2025/04/IMRJR.2025.020408.pdf>
 12. Tusconi, M., & Dursun, S. M. (2025). Violence and mental health. Focus on schizophrenia spectrum and psychotic disorders. *Frontiers in Psychiatry*.
<https://www.frontiersin.org/journals/psychiatry/articles/10.3389/fpsyt.2025.1618000/abstract>
 13. Zhong, X., She, J., & Wu, X. (2024). Tech for social good: Artificial intelligence and workplace safety. *Technology in Society*. <https://doi.org/10.1016/j.techsoc.2024.102745>
-

UDC 004 + 65.01

SCOPUS CODE 2213

<https://doi.org/10.36073/1512-0996-2025-3-309-316>

A New Stage of Occupational Safety: Integration of Artificial Intelligence in the Work Environment

- Nikoloz Shaverdashvili** Department of Engineering Safty and Emerngency Management, Georgian Technical University, Georgia, 0160, Tbilisi, 75, M. Kostava str.
E-mail: nikoloz298@gmail.com
- Nana Razmadze** Department of Engineering Safty and Emerngency Management, Georgian Technical University, Georgia, 0160, Tbilisi, 75, M. Kostava str.
E-mail: n.razmadze@gtu.ge
- Nino Ratiani** Department of Engineering Safty and Emerngency Management, Georgian Technical University, Georgia, 0160, Tbilisi, 75, M. Kostava str.
E-mail: n.ratiani@gtu.ge

Reviewers:

T. Kunchulia, Prof., Faculty of Mining and Geology, GTU

E-mail: t.kunchulia@gtu.ge

D. Butskhrikidze, Associate Professor, Faculty of Transport Systems and Mechanical Engineering, GTU

E-mail: d.butskhrikidze@gtu.ge

Abstract. This article discusses the role of Artificial Intelligence (AI) in improving occupational safety for employees in companies. It explores modern technological approaches, risk prediction methods, real-time monitoring systems, and digital training models. Examples from leading international companies are provided. The study demonstrates that AI serves not only as a technological innovation but also as a strategic management tool.

Keywords:

Artificial intelligence; Data analysis; Occupational safety; Risk management; Technological advancement.

განხილვის თარიღი 18.06.2025

შემოსვლის თარიღი 24.06.2025

ხელმოწერილია დასაბეჭდად 25.09.2025