

**ხელოვნური ინტელექტი და თანამედროვე ომი:
სტრატეგიული რღვევა, ეთიკური დილემები (საქართველოს შემთხვევა)**

ალიკა გუჩუა

პოლიტიკის მეცნიერების დოქტორი,
კავკასიის საერთაშორისო უნივერსიტეტის
ასოცირებული პროფესორი

ელ-ფოსტა: alika_guchua@ciu.edu.ge

ORCID: <https://orcid.org/0000-0003-0347-9574>

თორნიკე ზედელაშვილი

პოლიტიკის მეცნიერების დოქტორი, იუსტიციის
სამინისტროს ციფრული მმართველობის სააგენტოს
ინფორმაციული უსაფრთხოების დეპარტამენტის
თანამშრომელი; ინტერნეტგამოცემა „ლიდერის“

დამფუძნებელი; კავკასიის საერთაშორისო
უნივერსიტეტის მოწვეული ლექტორი

ელ-ფოსტა: thomaszedelashvili@gmail.com

ORCID: <https://orcid.org/0000-0003-2630-1779>

აბსტრაქტი. სტატია იკვლევს ხელოვნური ინტელექტის ტრანსფორმაციულ გავლენას თანამედროვე ომის ბუნებაზე, ქცევასა და შედეგებზე. სტატიის მიზანია წარმოადგინოს ხელოვნური ინტელექტის როლი თანამედროვე ომში და მასთან დაკავშირებული უსაფრთხოების ფაქტორები. მოცემული კვლევის მთავარი კვლევითი კითხვაა: როგორ ცვლის ხელოვნური ინტელექტის ტექნოლოგიების ინტეგრაცია ომის სტრატეგიულ, ეთიკურ და სამართლებრივ განზომილებებს და რა შედეგები მოჰყვება ამას საერთაშორისო უსაფრთხოებისა და ჰუმანიტარული სამართლისთვის? სტატია აანალიზებს ძირითად კვლევით კითხვებს სეკურიტიზაციის თეორიის, პოლიტიკური რეალიზმისა და ჰიბრიდული ომის გამოყენებით. კვლევა იყენებს პოლიტიკის ანალიზს, შინაარსობრივ ანალიზს და შემთხვევის შესწავლის მეთოდებს. ტექნოლოგიური დეტერმინიზმის, კონსეკვენციალისტური ეთიკისა და სამართლიანი ომის თეორიის თეორიულ ჩარჩოებზე დაყრდნობით, კვლევა იკვლევს ავტონომიური იარაღის სისტემების განლაგებას, ხელოვნური ინტელექტით გაძლიერებულ სამხედრო დაზვერვას და კიბერომის შესაძლებლობებს. დეტალური შემთხვევების შესწავლის გზით, მათ შორის რუსეთ-უკრაინის კონფლიქტის, ისრაელის „რკინის გუმბათის“ და ნატოს „DIANA“ ინიციატივის ჩათვლით, სტატია აჩვენებს, თუ როგორ ცვლის ხელოვნური ინტელექტი გადაწყვეტილების მიღების პროცესებს, აჩქარებს შეიარაღების რბოლას და ეჭვქვეშ აყენებს არსებულ სამართლებრივ ნორმებს. ნაშრომში განხილულია და გაანალიზებულია ხელოვნური ინტელექტისა და საქართველოს ეროვნული უსაფრთხოების პოლიტიკის საკითხი. დასკვნები ხაზს უსვამს გლობალური მარეგულირებელი ჩარჩოების, ეთიკური ზედამხედველობისა და მრავალმხრივი თანამშრომლობის აუცილებლობას, რათა უზრუნველყოფილი იყოს ხელოვნური ინტელექტის პასუხისმგებლიანი გამოყენება სამხედრო კონტექსტში.

საკვანძო სიტყვები: ხელოვნური ინტელექტი, ავტონომიური იარაღი, კიბერომი, სამხედრო სტრატეგია, შეიარაღების კონტროლი, საერთაშორისო ჰუმანიტარული სამართალი, სამართლიანი ომის თეორია, საქართველო.

შესავალი. ხელოვნური ინტელექტი სწრაფად წარმოიშვა, როგორც გლობალური უსაფრთხოების დესტრუქციული ძალა, რომელმაც შეცვალა ომის სტრატეგიული ლანდშაფტი და ხელახლა განსაზღვრა სამხედრო ძალაუფლების საზღვრები. ავტონომიური დრონებიდან, რომლებსაც აქვთ სასიკვდილო სიზუსტის უნარი, მტრის გადაადგილების პროგნოზირებად ალგორითმებამდე, ხელოვნური ინტელექტის ტექნოლოგიები სულ უფრო მეტად ინტეგრირდება თავდაცვის ინფრასტრუქტურასა და ოპერატიულ დოქტრინებში. ეს ტრანსფორმაცია არ არის მხოლოდ ტექნოლოგიური - ის ღრმად პოლიტიკური, ეთიკური და იურიდიულია. ხელოვნური ინტელექტის ინტეგრაცია ომში ფუნდამენტურ კითხვებს ბადებს ანგარიშვალდებულების, პროპორციულობისა და საერთაშორისო ჰუმანიტარული სამართლის მომავლის შესახებ.

აღნიშნული სტატიის ცენტრალური კვლევითი კითხვაა: როგორ ცვლის ხელოვნური ინტელექტის ტექნოლოგიების ინტეგრაცია ომის სტრატეგიულ, ეთიკურ და იურიდიულ განზომილებებს და რა შედეგები მოჰყვება ამას საერთაშორისო უსაფრთხოებისა და ჰუმანიტარული სამართლისთვის? ამ კითხვაზე პასუხის გასაცემად, კვლევა ეყრდნობა სამ ურთიერთდაკავშირებულ თეორიულ ჩარჩოს. პირველი, ტექნოლოგიური დეტერმინიზმი ამტკიცებს, რომ ტექნოლოგიური ინოვაცია იწვევს სოციალურ და პოლიტიკურ ცვლილებებს, ხშირად ადამიანის მოქმედებისგან დამოუკიდებლად. ომის კონტექსტში, ეს თეორია ეხმარება ახსნას, თუ როგორ ცვლის ხელოვნური ინტელექტის სისტემები სამხედრო დოქტრინებს, ოპერატიულ ქცევას და სტრატეგიული გადაწყვეტილების მიღებას (Altmann. & Sauer. 2017). მეორეც, კონსეკვენციალისტური ეთიკა აფასებს ქმედებებს მათი შედეგების მიხედვით, რაც ნორმატიულ ლინზას იძლევა საბრძოლო ხელოვნებაში ხელოვნური ინტელექტით გათვალისწინებული გადაწყვეტილებების მორალის შესაფასებლად, განსაკუთრებით მაშინ, როდესაც ადამიანის ზედამხედველობა შეზღუდულია ან საერთოდ არ არსებობს (Bode. & Huelss. 2022). მესამე, სამართლიანი ომის თეორია გვთავაზობს კლასიკურ ჩარჩოს ომის ლეგიტიმურობის (jus ad bellum) და ომის დროს ქცევის (jus in bello) შესაფასებლად, რომლებიც ორივეს ავტონომიური სისტემებისა და ალგორითმული დამიზნების განლაგებით ეჭვქვეშ აყენებს (Bode. & Huelss. 2022).

მეთოდები. მეთოდოლოგიურად, სტატია იყენებს თვისებრივ ანალიზს და თემატურ სინთეზს, ეყრდნობა რა ბოლოდროინდელ აკადემიურ ლიტერატურას, პოლიტიკურ დოკუმენტებს და რეალური სამყაროს შემთხვევების კვლევებს. მიზანია უზრუნველყოს ყოვლისმომცველი გაგება იმის შესახებ, თუ როგორ გარდაქმნის ხელოვნური ინტელექტი ომს და გამოავლინოს პოლიტიკური რეაგირება, რომელიც ამცირებს მის რისკებს. სტატია სტრუქტურირებულია ხუთ ძირითად ნაწილად: 1. ომის სტრატეგიული ტრანსფორმაცია ხელოვნური ინტელექტის მეშვეობით; 2. ავტონომიური იარაღის მიერ წარმოქმნილი ეთიკური დილემები; 3. ხელოვნური ინტელექტის როლი სამხედრო დაზვერვასა და დამიზნებაში; 4. ხელოვნური ინტელექტით გაძლიერებული კიბერომი და მისი შედეგები; და 5. გლობალური შეიარაღების რბოლა და მარეგულირებელი ხარვეზები. თითოეული სექცია აერთიანებს ემპირიულ მტკიცებულებებსა და თეორიულ რეფლექსიას, რათა შეიქმნას ხელოვნური ინტელექტის ომზე გავლენის ნიუანსირებული ანალიზი.

ხელოვნური ინტელექტი და ომის ტრანსფორმაცია

ხელოვნური ინტელექტის სამხედრო სისტემებში ინტეგრაცია ომის ბუნებისა და წარმოების პარადიგმის ცვლილებას აღნიშნავს. ისტორიულად, ტექნოლოგიურმა მიღწევებმა, როგორიცაა დენთი, მექანიზებული მანქანები და ბირთვული იარაღი, ხელახლა განსაზღვრა საბრძოლო და სტრატეგიული აზროვნება. ხელოვნური ინტელექტი აგრძელებს ამ ტრაექტორიას, რაც მანქანებს საშუალებას აძლევს შეასრულონ ტრადიციულად ადამიანებისთვის განკუთვნილი ამოცანები - თვალთვალი, საფრთხის შეფასება, გადაწყვეტილების მიღება და თუნდაც სასიკვდილო მოქმე-

დება. ეს ცვლილება არ არის მხოლოდ ოპერატიული; ის ონტოლოგიურია, რაც ცვლის ომის მნიშვნელობას და მასში ადამიანური აგენტების როლს.

ხელოვნური ინტელექტის ორმაგი დანიშნულების ბუნება ართულებს მის რეგულირებასა და ეთიკურ ზედამხედველობას. სამოქალაქო გამოყენებისთვის შემუშავებული ტექნოლოგიები - როგორცაა სახის ამოცნობა, ბუნებრივი ენის დამუშავება და ავტონომიური ნავიგაცია - სწრაფად ადაპტირდება სამხედრო გამოყენებისთვის. მაგალითად, კომერციული გამოსახულების ამოცნობის პროგრამული უზრუნველყოფა ხელახლა გამოიყენება ბრძოლის ველის იდენტიფიცირებისა და დამიზნებისთვის, ხოლო ავტონომიური მართვის ალგორითმები ხელს უწყობს უპილოტო სახმელეთო მანქანების განვითარებას (Cave. & Dignum. 2019). ეს კონვერგენცია აჩქარებს ხელოვნური ინტელექტის მილიტარიზაციას და იწვევს შეშფოთებას გამჭვირვალობის, ანგარიშვალდებულების და სამოქალაქო-სამხედრო საზღვრების ეროვნული შესახებ.

სტრატეგიულად, ხელოვნური ინტელექტი გთავაზობთ რამდენიმე უპირატესობას, რომელიც მიმზიდველია სამხედრო დამგეგმავებისთვის. ის აუმჯობესებს ოპერაციულ სიჩქარეს, სიზუსტეს და მასშტაბირებას, რაც რეალურ დროში მონაცემთა დამუშავებისა და სწრაფი გადაწყვეტილების მიღების საშუალებას იძლევა. 2022 წლის რუსეთ-უკრაინის კონფლიქტის დროს, ხელოვნური ინტელექტით გაძლიერებულმა დამიზნების სისტემებმა უკრაინის ძალებს საშუალება მისცა, თანამგზავრული სურათების, ღია წყაროებიდან მიღებული ინტელექტისა და მანქანური სწავლების ალგორითმების გამოყენებით, რუსული სამიზნეები უპრეცედენტო სიზუსტით ამოეცნოთ და დაეჯახათ (Geist. & Lohn. 2018). ანალოგიურად, ისრაელის ხელოვნური ინტელექტით გაძლიერებულმა „რკინის გუმბათის“ სისტემამ 2021 წლის დაზას კონფლიქტის დროს ათასობით რაკეტა ჩაჭრა, რითაც აჩვენა, თუ როგორ შეუძლია მანქანურ სწავლებას საფრთხის აღმოჩენისა და ჩაჭრის ოპტიმიზაცია (Gilli. & Gilli. 2021).

ჩინეთის სამხედრო დოქტრინა ხაზს უსვამს „ინტელექტუალურ ომს“, ხელოვნური ინტელექტის ინტეგრირებას სამეთაურო სისტემებში, ლოჯისტიკასა და საბრძოლო ველის სიმულაციებში. სახალხო განმათავისუფლებელმა არმიამ (PLA) დიდი ინვესტიცია ჩადო ხელოვნური ინტელექტის კვლევაში, მათ შორის ავტონომიურ წყალქვეშა მანქანებში, გუნდურ რობოტიკასა და ხელოვნურ ინტელექტზე დაფუძნებულ გადაწყვეტილების მხარდაჭერის სისტემებში. Gilli და Gilli (2021) თანახმად, ჩინეთი ხელოვნურ ინტელექტს სტრატეგიულ ხელშემწყობ ფაქტორად მიიჩნევს, რომელსაც შეუძლია ტრადიციული სამხედრო ნაკლოვანებების კომპენსირება და ძალთა გლობალური ბალანსის შეცვლა (Gilli. & Gilli. 2021).

თუმცა, ეს შესაძლებლობები ასევე მნიშვნელოვან რისკებს შეიცავს. გადაწყვეტილების მიღების უფრო სწრაფმა ციკლებმა შეიძლება შეამციროს დიპლომატიისა და კონფლიქტის დეესკალაციის შესაძლებლობები. ხელოვნური ინტელექტით დაფუძნებულ სისტემებს შეუძლიათ გამოიწვიონ ავტომატური რეაგირება აღქმულ საფრთხეებზე, რაც ზრდის გაუთვალისწინებელი ჩართულობისა და ესკალაციის ალბათობას. გარდა ამისა, ხელოვნური ინტელექტით გადაწყვეტილების მიღების ბუნდოვანება - რომელსაც ხშირად „შავი ყუთის“ პრობლემას უწოდებენ - ართულებს ანგარიშვალდებულებას და სტრატეგიულ სიგნალიზაციას. თუ მოწინააღმდეგეები ვერ გაიგებენ ან ვერ იწინასწარმეტყველებენ ხელოვნური ინტელექტით მართულ ქცევას, იზრდება არასწორი გაანგარიშებისა და შემთხვევითი კონფლიქტის რისკი (Geist. & Lohn. 2018).

ხელოვნური ინტელექტის მეშვეობით ომის ტრანსფორმაცია ასევე გამოწვევას უქმნის ტრადიციულ სამეთაურო სტრუქტურებსა და ოპერატიულ დოქტრინებს. მეთაურ პიროვნებებს ახლა უწევთ ურთიერთქმედება რთულ ალგორითმულ სისტემებთან, ხშირად ეყრდნობიან მანქანების მიერ გენერირებულ რეკომენდაციებს სიცოცხლისა და სიკვდილის გადაწყვეტილებებისთვის. ეს ცვლილება კითხვებს ბაღებს ადამიანის განსჯის ეროვნული, მორალური პასუხისმგებლობის დელეგი-

რების და ჯარისკაცებსა და გადაწყვეტილების მიმღებ პირებზე ფსიქოლოგიური ზემოქმედების შესახებ. როგორც Scharre (2018) ამტკიცებს, „კენტავრი მეომრების“ - ადამიანი-მანქანის ჰიბრიდების - აღზვევა მოითხოვს წვრთნის, ლიდერობისა და ეთიკური რეფლექსიის ახალ მოდელებს.

მოკლედ, ხელოვნური ინტელექტი არ არის მხოლოდ ომის ინსტრუმენტი; ეს არის ძალა, რომელიც ცვლის სამხედრო ძალაუფლების სტრატეგიულ, ეთიკურ და ინსტიტუციურ საფუძვლებს. მისი ომში ინტეგრაცია მოითხოვს დოქტრინების, სამართლებრივი ნორმებისა და ადამიან-მანქანის ურთიერთობების გადახედვას. შემდეგი ნაწილი განიხილავს ამ ტრანსფორმაციის ერთ-ერთ ყველაზე საკამათო ასპექტს: ავტონომიური იარაღის განლაგებას და მათ მიერ წარმოქმნილ ეთიკურ დილემებს.

ავტონომიური იარაღი და ეთიკური დილემები

ხელოვნური ინტელექტის სამხედრო კონტექსტში გამოყენების ყველაზე საკამათო საკითხებს შორისაა ლეტალური ავტონომიური იარაღის სისტემების (LAWS) შემუშავება და განლაგება. ამ სისტემებს შეუძლიათ სამიზნეების შერჩევა და მათზე თავდასხმა ადამიანის პირდაპირი ჩარევის გარეშე, რაც ღრმა ეთიკურ, სამართლებრივ და სტრატეგიულ შეშფოთებას იწვევს. LAWS-ის გარშემო დებატები არ არის მხოლოდ ტექნიკური - ის ფუნდამენტურად ფილოსოფიურია და ეხება ომში ავტომატიზაციის მოქმედების ბუნებას, პასუხისმგებლობას და მორალურ საზღვრებს.

ავტონომიური იარაღი ეჭვქვეშ აყენებს საერთაშორისო ჰუმანიტარული სამართლის ფუნდამენტურ პრინციპებს, განსაკუთრებით განსხვავების, პროპორციულობისა და ანგარიშვალდებულების მოთხოვნებს. განსხვავების პრინციპი მოითხოვს, რომ მეზრძოლებმა განასხვავონ სამხედრო სამიზნეები და მშვიდობიანი მოქალაქეები, ხოლო პროპორციულობა მოითხოვს, რომ თავდასხმის მოსალოდნელი სამხედრო უპირატესობა აღემატებოდეს მშვიდობიანი მოსახლეობისთვის პოტენციურ ზიანს. ადამიანის განსჯის არარსებობის შემთხვევაში, გაურკვეველია, თუ როგორ შეუძლიათ ავტონომიურ სისტემებს საიმედოდ დაიცვან ეს პრინციპები. როგორც Cave და Dignum (2019) ამტკიცებენ, ლეტალური გადაწყვეტილებების მიღების მანქანებისთვის დელეგირება ომის ეთიკური არქიტექტურის ძირს უთხრის (Cave. & Dignum. 2019).

რეალური სამყაროს მოვლენები ასახავს ამ შეშფოთების აქტუალურობას. 2020 წელს გავრცელდა ინფორმაცია, რომ თურქეთის STM Kargu-2 დრონი შესაძლოა ლიბიაში შეტაკების დროს ავტონომიურად მოქმედებდა, რაც წარმოადგენს ხელოვნური ინტელექტით მართული სასიკვდილო ძალის გამოყენების ერთ-ერთ პირველ ცნობილ შემთხვევას ადამიანის ზედამხედველობის გარეშე (Saylor. 2020). მიუხედავად იმისა, რომ დეტალები სადავოა, ინციდენტმა გამოიწვია საერთაშორისო დებატები ავტონომიის ზღურბლისა და რეგულირების საჭიროების შესახებ. ანალოგიურად, შეერთებულმა შტატებმა გამოსცა ავტონომიური სისტემები, რომლებსაც შეუძლიათ გუნდურად ქცევის განხორციელება, სადაც მრავალი დრონი კოორდინაციას უწევს თავდასხმებს ცენტრალიზებული კონტროლის გარეშე. ეს შესაძლებლობები გვთავაზობს ტაქტიკურ უპირატესობებს, მაგრამ ასევე ბადებს კითხვებს საგანგებო ქცევისა და საბრძოლო სცენარებში არაპროგნოზირებადობის შესახებ.

ეთიკური თვალსაზრისით, კონსეკვენციალისტური მსჯელობა იძლევა ერთ ლინზას ავტონომიური იარაღის გამოყენების შესაფასებლად. თუ ასეთ სისტემებს შეუძლიათ შეამცირონ მშვიდობიან მოსახლეობაში მსხვერპლი, მინიმუმამდე დაიყვანონ ადამიანური შეცდომები და გააძლიერონ ოპერატიული ეფექტურობა, ისინი შეიძლება ჩაითვალოს მორალურად დასაშვებად. თუმცა, ეს ლოგიკა ვარაუდობს, რომ ხელოვნური ინტელექტის სისტემებს შეუძლიათ საიმედოდ ინტერპრეტაცია გაუწიონ რთულ საბრძოლო გარემოს, რაც გარანტირებულისგან შორს არის. მანქანური სწავლების ალგორითმები მხოლოდ იმდენად კარგია, რამდენადაც მათი სასწავლო მონაცემები და დინამიურ, მაღალი რისკის მქონე გარემოში არასწორი კლასიფიკაციის ან გაუთვალისწინებელი ექსტრემულობის რისკი მნიშვნელოვანია (Altmann. & Sauer. 2017).

გარდა ამისა, პასუხისმგებლობის საკითხი კვლავ გადაუჭრელი რჩება. ტრადიციულ ომში, უკანონო ქმედებების პასუხისმგებლობა შეიძლება ადამიანებს მოქმედ პირებს - მეთაურებს, ჯარისკაცებს ან პოლიტიკურ ლიდერებს მივაკუთვნოთ. ავტონომიური სისტემების შემთხვევაში, პასუხისმგებლობის ჯაჭვი ბუნდოვანი ხდება. თუ დრონი ომის დანაშაულს ჩაიდენს, ვინ არის პასუხისმგებელი? პროგრამისტი, რომელმაც კოდი დაწერა? მეთაური, რომელმაც სისტემა განალაგა? მწარმოებელი, რომელმაც აპარატურა შექმნა? იურიდიული მიმართულებით მეცნიერებმა შემოგვთავაზეს სხვადასხვა მოდელი, მათ შორის მკაცრი პასუხისმგებლობის რეჟიმები და სავალდებულო „ადამიანის ციკლში ჩართვის“ პროტოკოლები, მაგრამ კონსენსუსი კვლავ გაურკვეველია (Scharre. 2018).

სამართლიანი ომის თეორია დამატებით წარმოდგენას გვთავაზობს ავტონომიური იარაღის მიერ წარმოქმნილ ეთიკურ დილემებზე. მაგალითად, სწორი განზრახვის პრინციპი მოითხოვს, რომ ომი მორალურად გამართლებული მიზეზების გამო დაიწყოს. თუ ავტონომიური სისტემები გამოიყენება ალგორითმული პროგნოზების საფუძველზე პრევენციული დარტყმების განსახორციელებლად, განზრახვის მორალური სიცხადე ბუნდოვანი ხდება. ანალოგიურად, უკანასკნელი დონისძიების პრინციპი ვარაუდობს, რომ ადამიანმა მოქმედ პირებმა ძალადობამდე ამოწურეს დიპლომატიური ვარიანტები. სამყაროში, სადაც მანქანები გადაწყვეტილებებს მანქანების სიჩქარით იღებენ, დიპლომატიის დროითი სივრცე შეიძლება შემცირდეს ან საერთოდ აღმოიფხვრას (Bode. & Huelss 2022).

ავტონომიური იარაღის რეგულირების მცდელობებმა იმპულსი მოიპოვა, მაგრამ ფრაგმენტული რჩება. გარკვეული ჩვეულებრივი იარაღის შესახებ გაეროს კონვენცია (CCW) 2014 წლიდან იწვევდა ექსპერტთა შეხვედრებს კანონების განსახილველად, თუმცა არანაირი სავალდებულო ხელშეკრულება არ შექმნილა. ზოგიერთი სახელმწიფო, მათ შორის ავსტრია და ბრაზილია, ეთიკურ და ჰუმანიტარულ საკითხებზე დაყრდნობით პრევენციულ აკრძალვას უჭერს მხარს. სხვები, როგორიცაა შეერთებული შტატები და რუსეთი, ეწინააღმდეგებიან რეგულაციას და ამტკიცებენ, რომ ავტონომია აძლიერებს სამხედრო ეფექტურობას და შეკავების ძალას (Boulainin. & Verbruggen. 2017).

სამოქალაქო საზოგადოების ორგანიზაციებიც მობილიზებული არიან ავტონომიური იარაღის წინააღმდეგ. 2013 წელს დაწყებული მკვლელი რობოტების შეჩერების კამპანია მოუწოდებს საერთაშორისო ხელშეკრულებისკენ, რომელიც კრძალავს სრულად ავტონომიურ სასიკვდილო სისტემებს. მათმა მხარდაჭერამ გავლენა მოახდინა საზოგადოებრივ დისკურსზე და ეროვნული მთავრობები აიძულა, განეხილათ ეთიკური სახელმძღვანელო პრინციპები, მაგრამ სამართლებრივი კოდიფიცირება ჯერ კიდევ შორს არის.

მოკლედ, ავტონომიური იარაღი წარმოადგენს კრიტიკულ ეტაპს ომის ევოლუციაში. ისინი გვთავაზობენ სტრატეგიულ უპირატესობებს, მაგრამ წარმოადგენენ მნიშვნელოვან ეთიკურ და სამართლებრივ რისკებს. მკაფიო ანგარიშვალდებულების მექანიზმების არარსებობა, გაუთვალისწინებელი ესკალაციის პოტენციალი და ჰუმანიტარული ნორმების ეროვნული სასწრაფო ყურადღებას მოითხოვს. რადგან ხელოვნური ინტელექტი აგრძელებს განვითარებას, საერთაშორისო საზოგადოებამ უნდა გაუმკლავდეს მანქანებისთვის სასიკვდილო ძალის დელეგირების მორალურ და სამართლებრივ შედეგებს. შემდეგი განყოფილება იკვლევს, თუ როგორ ცვლის ხელოვნური ინტელექტი სამხედრო დაზვერვასა და დამიზნებას, რაც კიდევ უფრო ართულებს თანამედროვე კონფლიქტის ეთიკურ რელიეფს.

ხელოვნური ინტელექტი სამხედრო დაზვერვასა და დამიზნებაში

ხელოვნურმა ინტელექტმა რევოლუცია მოახდინა სამხედრო დაზვერვის სფეროში, შემოიღო შესაძლებლობები, რომლებიც გაცილებით აღემატება ტრადიციულ ადამიანურ ანალიზს. მანქანური სწავლების, ბუნებრივი ენის დამუშავებისა და კომპიუტერული ხედვის მეშვეობით, ხელოვნური

ინტელექტის სისტემებს შეუძლიათ რეალურ დროში მონაცემების უზარმაზარი მოცულობის დამუშავება, ნიმუშების იდენტიფიცირება და ქმედითი ინფორმაციის გენერირება. ეს ტექნოლოგიები სულ უფრო ხშირად გამოიყენება სტრატეგიული დაგეგმვის, ბრძოლის ველზე ცნობიერების ამაღლებისა და ზუსტი დამიზნების მხარდასაჭერად, რაც ცვლის სამხედრო გადაწყვეტილებების მიღებისა და შესრულების წესს.

ხელოვნური ინტელექტის ერთ-ერთი ყველაზე მნიშვნელოვანი უპირატესობა სამხედრო დაზვერვაში არის მისი უნარი, სინთეზირება გაუკეთოს ჰეტეროგენულ მონაცემთა წყაროებს. თანამგზავრული გამოსახულებები, დრონების ჩანაწერები, ჩაჭრილი კომუნიკაციები და სოციალური მედიის კონტენტი შეიძლება ინტეგრირებული იყოს გაერთიანებულ პლატფორმებში, რომლებიც მეთაურებს ყოვლისმომცველ სიტუაციურ ცნობიერებას ანიჭებენ. მაგალითად, რუსეთ-უკრაინის კონფლიქტის დროს, უკრაინის ძალებმა გამოიყენეს ხელოვნური ინტელექტით გაუმჯობესებული სისტემები თანამგზავრული მონაცემებისა და ღია წყაროს დაზვერვის გასაანალიზებლად, რაც საშუალებას აძლევდა რუსული ჯარების გადაადგილებისა და არტილერიის პოზიციების სწრაფ იდენტიფიცირებას (King. 2024). ეს სისტემები იყენებდნენ კონვოლუციურ ნეირონულ ქსელებს რელიეფსა და ინფრასტრუქტურაში ანომალიების აღმოსაჩენად, რაც საშუალებას იძლეოდა უფრო ზუსტი დამიზნებისა და თანმხლები ზიანის შემცირების.

შეერთებულმა შტატებმა ასევე დიდი ინვესტიციები ჩადო ხელოვნურ ინტელექტზე დაფუძნებულ დაზვერვის პლატფორმებში. 2017 წელს თავდაცვის დეპარტამენტის მიერ დაწყებული პროექტი Maven მიზნად ისახავდა დრონების მიერ გადაღებული ვიდეოჩანაწერების ანალიზის ავტომატიზაციას კომპიუტერული ხედვის ალგორითმების გამოყენებით. მიუხედავად იმისა, რომ პროექტს Google-ის თანამშრომლების მხრიდან ეთიკური უკმაყოფილება მოჰყვა, მან აჩვენა ხელოვნური ინტელექტის პოტენციალი, დააჩქაროს სადაზვერვო ციკლები და შეამციროს კოგნიტური დატვირთვა ანალიტიკოსებზე (Geist. & Lohn. 2018). ცოტა ხნის წინ, უკრაინის უსაფრთხოების დახმარების ჯგუფმა ხელოვნური ინტელექტი გამოიყენა საბრძოლო ველის მონაცემების პროგნოზირებად ანალიტიკასთან შესარწყმელად, რაც აძლიერებს ოპერატიულ კოორდინაციას და რესურსების განაწილებას.

პროგნოზირებადი დამიზნება კიდევ ერთი სფეროა, სადაც ხელოვნურმა ინტელექტმა მნიშვნელოვანი წინსვლა განიცადა. ისტორიული ნიმუშების, ქვევითი ინდიკატორების და გარემოსდაცვითი ცვლადების ანალიზით, ხელოვნური ინტელექტის სისტემებს შეუძლიათ მტრის ქმედებების პროგნოზირება და პრევენციული რეაგირების რეკომენდაცია. ეს შესაძლებლობები განსაკუთრებით ღირებულია კონტრ-ამბოხებულებებსა და ურბანულ ომებში, სადაც მოწინააღმდეგეები ხშირად ერწყმიან მშვიდობიან მოსახლეობას. თუმცა, პროგნოზირებადი დამიზნება ეთიკურ შეშფოთებას იწვევს პროფილირების, მიკერძოების და ცრუ დადებითი შედეგების რისკის შესახებ. არასრულ ან დამახინჯებულ მონაცემთა ნაკრებებზე გაწვრთნილმა ალგორითმებმა შეიძლება არასწორად ამოიცნონ პირები ან არასწორად განმარტონ განზრახვა, რაც გამოიწვევს არასწორ დამიზნებას და საერთაშორისო სამართლის პოტენციურ დარღვევებს (Taddeo. & Floridi. 2018).

ხელოვნური ინტელექტის გადაწყვეტილების მიღების ბუნდოვანება კიდევ უფრო ართულებს ანგარიშვალდებულებას. მანქანური სწავლების მრავალი მოდელი მოქმედებს როგორც „შავი ყუთი“, რაც შედეგებს იძლევა გამჭვირვალე მსჯელობის გარეშე. ინტერპრეტაციის ეს ნაკლებობა სირთულეებს უქმნის სამხედრო ზედამხედველობას სამართლებრივი განხილვისთვის. თუ სამიზნედ განსაზღვრის გადაწყვეტილება გაუმჭირვალე ალგორითმს ეფუძნება, რთული ხდება მისი შესაბამისობის შეფასება აუცილებლობისა და პროპორციულობის პრინციპებთან. როგორც Altmann და Sauer (2017) აღნიშნავენ, ხელოვნური ინტელექტის ინტეგრაცია სამიზნედ განსაზღვრის სისტემებში განმარტებისა და აუდიტის ახალ სტანდარტებს მოითხოვს (Altmann. & Sauer. 2017).

გარდა ამისა, მოწინააღმდეგეები სულ უფრო მეტად აცნობიერებენ ხელოვნური ინტელექტის როლს დაზვერვაში და ავითარებენ საპასუხო ზომებს. მონაცემთა „მოწამვლა“, შეწინააღმდეგების მიერ შეყვანილი მონაცემები და გაყალბების ტექნიკა შეიძლება ხელოვნური ინტელექტის სისტემებს შეცდომაში შეჰყავდეს, რაც არასწორ კლასიფიკაციას და ოპერაციულ ჩავარდნებს იწვევს. 2023 წელს ნატოს DIANA ინიციატივამ დაიწყო კვლევითი პროგრამების სერია, რომელიც მიზნად ისახავს ხელოვნური ინტელექტის სისტემების მდგრადობის გაზრდას ასეთი თავდასხმების მიმართ. ეს ძალისხმევა მოიცავს შეწინააღმდეგების მიერ სწავლების პროტოკოლების და ანომალიების აღმოჩენის ალგორითმების შემუშავებას, რომლებიც შექმნილია მანიპულირების მცდელობების რეალურ დროში აღმოსაჩენად (Kaspersen. 2021).

მიზნობრივად ხელოვნური ინტელექტის ეთიკური შედეგები ბრძოლის ველზე ეფექტურობის ფარგლებს სცდება. ისინი ეხება ადამიანის მოქმედების ფუნდამენტურ საკითხს სიცოცხლისა და სიკვდილის გადაწყვეტილებებში. მიუხედავად იმისა, რომ ხელოვნურ ინტელექტს შეუძლია გააძლიეროს ადამიანის განსჯა, არსებობს რისკი, რომ მეთაურები შეიძლება ზედმეტად დაეყრდნონ ალგორითმულ რეკომენდაციებს, განსაკუთრებით დროის ზეწოლის ქვეშ. ამ ფენომენს, რომელიც ცნობილია როგორც ავტომატიზაციის მიკერძოება, შეუძლია კრიტიკული აზროვნების შერყევა და ანგარიშვალდებულების შემცირება. როგორც Scharre (2018) ხაზს უსვამს, მიზნობრივ გადაწყვეტილებებზე მნიშვნელოვანი ადამიანის კონტროლის შენარჩუნება აუცილებელია მორალური პასუხისმგებლობის შესანარჩუნებლად და შეიარაღებული კონფლიქტის კანონების დასაცავად (Scharre. P. 2018).

ოპერაციულ საზრუნავებთან ერთად, ხელოვნური ინტელექტის გამოყენება დაზვერვაში სტრატეგიულ და გეოპოლიტიკურ საკითხებს წარმოშობს. განვითარებული ხელოვნური ინტელექტის შესაძლებლობების მქონე სახელმწიფოებმა შესაძლოა არაპროპორციული გავლენა მოიპოვონ გლობალურ საქმეებში, რაც ძალაუფლებისა და შეკავების ასიმეტრიას გამოიწვევს. ხელოვნური ინტელექტის ტექნოლოგიების გავრცელება ასევე ართულებს შეიარაღების კონტროლს, რადგან პროგრამულ უზრუნველყოფაზე დაფუძნებული სისტემების მონიტორინგი და გადამოწმება უფრო რთულია, ვიდრე ტრადიციული იარაღის. როგორც Boulanin და Verbruggen (2017) ამტკიცებენ, ხელოვნურ ინტელექტზე დაფუძნებული სამიზნე პლატფორმების გავრცელება მოითხოვს სახელმწიფოებს შორის გამჭვირვალობის, გადამოწმებისა და ნდობის აღდგენის ახალ ჩარჩოებს (Boulanin. & Verbruggen. 2017).

შეჯამების სახით, ხელოვნურმა ინტელექტმა შეცვალა სამხედრო დაზვერვა და მიზნობრივი მართვა, უზრუნველყო უპრეცედენტო შესაძლებლობები მონაცემთა ანალიზისთვის, პროგნოზირებისა და გადაწყვეტილების მხარდაჭერისთვის. თუმცა, ამ სარგებელს თან ახლავს მნიშვნელოვანი ეთიკური, იურიდიული და სტრატეგიული რისკები. ამ სფეროში ხელოვნური ინტელექტის პასუხისმგებლიანი გამოყენების უზრუნველყოფა მოითხოვს ძლიერ ზედამხედველობის მექანიზმებს, გამჭვირვალე ალგორითმებს და ადამიანის განსჯის გადამოწმებას კრიტიკული გადაწყვეტილების მიღებისას. შემდეგი ნაწილი განიხილავს, თუ როგორ ცვლის ხელოვნური ინტელექტი კიბერომს, კიდევ უფრო აფართოებს თანამედროვე კონფლიქტების მასშტაბებსა და სირთულეს.

კიბერომი და ხელოვნური ინტელექტით გაძლიერებული თავდაცვა

ხელოვნური ინტელექტი კიბერომის ცენტრალურ კომპონენტად იქცა, რაც უპრეცედენტო სიჩქარითა და დახვეწილობით უზრუნველყოფს როგორც შეტევით, ასევე თავდაცვით ოპერაციებს. ტრადიციული ომისგან განსხვავებით, კიბერკონფლიქტი ციფრულ სფეროში ვითარდება, სადაც ატრიბუციის დადგენა რთულია, საზღვრები ბუნდოვანია და სამოქალაქო ინფრასტრუქტურა ხშირად პირდაპირი სამიზნეა. ხელოვნური ინტელექტი აძლიერებს ამ სფეროს საფრთხის აღმოჩენის

ავტომატიზირებით, რეაგირების სტრატეგიების ოპტიმიზაციით და რეალურ დროში ადაპტური სწავლების ჩატარების გზით.

შეტევით, ხელოვნური ინტელექტის გამოყენება შესაძლებელია სისტემის დაუცველობის იდენტიფიცირებისთვის, პოლიმორფული მავნე პროგრამების გენერირებისთვის და მასშტაბური ფიშინგის კამპანიების ჩასატარებლად. ეს შესაძლებლობები საშუალებას აძლევს სახელმწიფო და არასახელმწიფო აქტორებს შეადგინონ უსაფრთხო ქსელებში, შეაფერხონ კომუნიკაციები და მანიპულირონ მონაცემები. მაგალითად, რუსეთსა და უკრაინას შორის ომის ადრეულ ეტაპზე, ხელოვნური ინტელექტით აღჭურვილი კიბერინსტრუმენტები, როგორც ამბობენ, გამოიყენებოდა უკრაინის მთავრობის ვებსაიტებსა და ენერგეტიკულ ინფრასტრუქტურაზე თავდასხმისთვის, რაც იწვევდა დროებით შეფერხებებს და ავრცელებდა დეზინფორმაციას (Kaspersen. 2021).

თავდაცვის სფეროში, ხელოვნური ინტელექტი მნიშვნელოვან როლს ასრულებს შეჭრის აღმოჩენის სისტემებში (IDS), ანომალიების ამოცნობასა და პროგრამული უზრუნველყოფის ავტომატურ განახლებებში. მანქანური სწავლების ალგორითმებს შეუძლიათ ქსელური ტრაფიკის ანალიზი, უჩვეულო ნიმუშების იდენტიფიცირება და მავნე კოდის იზოლირება მის გავრცელებამდე. ნატოს კიბერთავდაცვის კოოპერატიული ცენტრის მიერ ხაზგასმულია ხელოვნური ინტელექტის მნიშვნელობა მდგრადი კიბერთავდაცვის შექმნაში, განსაკუთრებით კი ისეთი კრიტიკული ინფრასტრუქტურის დაცვაში, როგორიცაა საავადმყოფოები, ელექტროქსელები და ფინანსური სისტემები (Bode. & Huelss. 2022).

თუმცა, ხელოვნური ინტელექტის გამოყენება კიბერომში სერიოზულ ეთიკურ და სამართლებრივ შეშფოთებას იწვევს. სამოქალაქო ინფრასტრუქტურაზე თავდასხმებმა შეიძლება დაარღვიოს საერთაშორისო ჰუმანიტარული სამართალი, განსაკუთრებით თუ ისინი ზიანს აყენებენ მშვიდობიან მოსახლეობას. პროპორციულობის პრინციპის გამოყენება კიბერსივრცეში რთულია, სადაც თავდასხმის შედეგები შეიძლება იყოს არაპროგნოზირებადი ან დაგვიანებული. უფრო მეტიც, კიბერშეტევისთვის მიკუთვნების პრობლემა ართულებს სამართლებრივ პასუხისმგებლობას და დიპლომატიურ რეაგირებას.

ხელოვნური ინტელექტი ასევე ქმნის ესკალაციის რისკს ავტომატური შურისძიების გზით. ზოგიერთი სამხედრო დოქტრინა ითვალისწინებს „აქტიური თავდაცვის“ სისტემებს, რომლებიც ავტონომიურად რეაგირებენ კიბერშეტევებზე. თუ ასეთი სისტემები არასწორად ამოიცნობენ საფრთხეს ან არაპროპორციულად უპასუხებენ, მათ შეუძლიათ გაუთვალისწინებელი კონფლიქტის პროვოცირება. გეისტი და ლოუნი (2018) აფრთხილებენ, რომ ხელოვნური ინტელექტით მართულ კიბეროპერაციებმა შეიძლება შეამციროს ადამიანებისთვის გადაწყვეტილების მიღების ფანჯარა, რაც გაზრდის არასწორი გაანგარიშების ალბათობას (Geist. & Lohn. 2018).

ამ რისკების მოსაგვარებლად, საერთაშორისო ორგანიზაციებმა შემოგვთავაზეს კიბერსივრცეში პასუხისმგებლიანი ქცევის ნორმები. გაეროს სამთავრობო ექსპერტთა ჯგუფმა (UNGGE) მოუწოდა გამჭვირვალობის, თავშეკავებისა და თანამშრომლობისკენ კიბეროპერაციებში. ევროკავშირის კიბერსოლიდარობის აქტი, რომელიც 2024 წელს იქნა მიღებული, მოიცავს ხელოვნურ ინტელექტზე დაფუძნებული საფრთხის აღმოჩენისა და საზღვრისპირა ინციდენტებზე რეაგირების დებულებებს. ეს ინიციატივები ასახავს ხელოვნური ინტელექტის როლის მზარდ აღიარებას კიბერთავდაცვაში და საერთო სტანდარტების საჭიროების შესახებ.

საერთო ჯამში, ხელოვნურმა ინტელექტმა გააფართოვა კიბერომის მასშტაბები და დახვეწილობა, რაც უზრუნველყოფს ძლიერ ინსტრუმენტებს როგორც თავდასხმისთვის, ასევე თავდაცვისთვის. მიუხედავად იმისა, რომ ეს ტექნოლოგიები აძლიერებს ოპერატიულ ეფექტურობას, ისინი ასევე ეჭვქვეშ აყენებენ სამართლებრივ ნორმებსა და ეთიკურ საზღვრებს. პასუხისმგებლიანი

გამოყენების უზრუნველყოფა მოითხოვს ძლიერ ზედამხედველობას, საერთაშორისო თანამშრომლობას და სამოქალაქო ინფრასტრუქტურის ციფრული ზიანისგან დაცვის ვალდებულებას.

გლობალური შეიარაღების რბოლა და პოლიტიკური ხარვეზები

ხელოვნური ინტელექტის სტრატეგიულმა ღირებულებამ გლობალური შეიარაღების რბოლა გამოიწვია, წამყვანი სახელმწიფოები დიდ ინვესტიციებს დებენ ავტონომიურ სისტემებში, ალგორითმულ ომსა და მონაცემებზე დაფუძნებულ სამხედრო პლატფორმებში. შეერთებულმა შტატებმა, ჩინეთმა, რუსეთმა და ევროკავშირმა შეიმუშავეს ეროვნული ხელოვნური ინტელექტის სტრატეგიები, რომლებიც მოიცავს სამხედრო კომპონენტებს. ეს ინიციატივები პრიორიტეტს ანიჭებს ავტონომიას, სიჩქარეს და ინფორმაციულ უპირატესობას, რაც ცვლის ძალთა ბალანსს საერთაშორისო ურთიერთობებში.

ჩინეთის „შემდეგი თაობის ხელოვნური ინტელექტის განვითარების გეგმა“ აშკარად აკავშირებს ხელოვნურ ინტელექტს ეროვნულ უსაფრთხოებასთან, ხაზს უსვამს ინტელექტუალურ მართვის სისტემებსა და უპილოტო საბრძოლო პლატფორმებს. შეერთებულმა შტატებმა უპასუხა ხელოვნური ინტელექტის ერთობლივი ცენტრის (JAIC) შექმნით, რომელიც კოორდინაციას უწევს ხელოვნური ინტელექტის ინტეგრაციას თავდაცვის სამინისტროში. რუსეთი ფოკუსირებულია ხელოვნურ ინტელექტზე დაფუძნებულ ელექტრონულ ომსა და ავტონომიურ სახმელეთო მანქანებზე, ხოლო ევროკავშირი ხაზს უსვამს ეთიკურ პრინციპებსა და ორმაგი დანიშნულების რეგულაციებს (Gilli. & Gilli. 2021).

ამ მოვლენების მიუხედავად, არ არსებობს სავალდებულო საერთაშორისო ხელშეკრულება, რომელიც არეგულირებს ხელოვნური ინტელექტის სამხედრო გამოყენებას. არსებული შეიარაღების კონტროლის შეთანხმებები, როგორიცაა ჟენევის კონვენციები და ევროპაში ჩვეულებრივი შეიარაღებული ძალების შესახებ ხელშეკრულება, არ ეხება ალგორითმულ სისტემებს ან ავტონომიურ პლატფორმებს. ახალ ჩარჩოებზე შეთანხმების მცდელობებს წინააღმდეგობა შეხვდა, რადგან სახელმწიფოები შიშობენ სტრატეგიული უპირატესობის დაკარგვის ან მგრძნობიარე ტექნოლოგიების გამჟღავნების (Boulanin. & Verbruggen. 2017).

რეგულაციების არარსებობა ქმნის ნაყოფიერ გარემოს ექსპერიმენტებისა და განხორციელებისთვის. ზედამხედველობის გარეშე, ხელოვნურმა ინტელექტმა შეიძლება შეამციროს კონფლიქტის ზღვარი, შესაძლებელი გახადოს პრევენციული დარტყმები და შეაფერხოს შეკავების ეფექტი. Altmann და Sauer (2017) ამტკიცებენ, რომ ავტონომიურმა სისტემებმა შეიძლება დესტაბილიზაცია მოახდინონ სტრატეგიული სტაბილურობის დესტაბილიზაციის გზით, გაურკვევლობის შემოღებით და სამხედრო ოპერაციების გამჭვირვალობის შემცირებით (Altmann. & Sauer. 2017).

სამოქალაქო საზოგადოების ორგანიზაციები მოითხოვენ საერთაშორისო ნორმებისა და ხელშეკრულებების მიღებას, რომლებიც არეგულირებენ ხელოვნური ინტელექტის გამოყენებას ომში. მკვლელი რობოტების შეჩერების კამპანია მხარს უჭერს სრულად ავტონომიური ლეტალური სისტემების აკრძალვას, ეთიკურ და ჰუმანიტარულ საკითხებზე დაყრდნობით. გაეროს მაღალი დონის პანელმა განვითარებადი ტექნოლოგიების შესახებ რეკომენდაცია გაუწია სამხედრო ხელოვნური ინტელექტის აპლიკაციების გლობალური რეესტრის შექმნას, მაგრამ მისი განხორციელება კვლავ გაურკვეველია.

პოლიტიკურმა რეაგირებამ უნდა გადაჭრას რამდენიმე ძირითადი საკითხი. პირველ რიგში, სახელმწიფოები უნდა შეთანხმდნენ ავტონომიის განმარტებებსა და ზღვრებზე. მეორეც, უნდა შემუშავდეს ვერიფიკაციის მექანიზმები შესაბამისობის უზრუნველსაყოფად. მესამე, ეთიკური სტანდარტები უნდა იყოს ჩადებული სისტემის დიზაინში, მათ შორის ადამიანური ზედამხედველო-

ბისა და ახსნა-განმარტების მოთხოვნების ჩათვლით. და ბოლოს, მრავალმხრივი თანამშრომლობა აუცილებელია ნდობის ასამაღლებლად და შეიარაღების რბოლის თავიდან ასაცილებლად.

შეჯამების სახით, ხელოვნური ინტელექტის ომში გლობალური შეიარაღების რბოლა მნიშვნელოვან რისკებს უქმნის საერთაშორისო უსაფრთხოებას. რეგულაციების გარეშე, ამ ტექნოლოგიებმა შეიძლება ძირი გამოუთხაროს ჰუმანიტარულ ნორმებს, დესტაბილიზაცია მოახდინოს შეკავების მექანიზმებზე და გაზარდოს კონფლიქტის ალბათობა. საერთაშორისო საზოგადოებამ დაუყოვნებლივ უნდა იმოქმედოს ისეთი ჩარჩოს შესაქმნელად, რომელიც უზრუნველყოფს სამხედრო ხელოვნური ინტელექტის პასუხისმგებლიან განვითარებასა და გამოყენებას.

ხელოვნური ინტელექტი და ეროვნული უსაფრთხოება საქართველოში

რადგან ხელოვნური ინტელექტი (AI) გლობალური უსაფრთხოების პარადიგმებს ტრანსფორმირებს, ისეთი პატარა სახელმწიფოები, როგორიც საქართველოა, ორმაგი გამოწვევის წინაშე დგანან: ახალი ტექნოლოგიების გამოყენება საკუთარი თავდაცვის მოდერნიზებისთვის და ამავედროულად ხელოვნური ინტელექტით აღჭურვილი ომის ეთიკურ, სამართლებრივ და გეოპოლიტიკურ სირთულეებში ნავიგაცია. აღმოსავლეთისა და დასავლეთის გზაჯვარედინზე მდებარე საქართველოს სტრატეგიული დაუცველობა, რომელიც გამწვავებულია გადაუჭრელი ტერიტორიული კონფლიქტებითა და რევიზიონისტულ ძალებთან სიახლოვით, ხელოვნური ინტელექტის ინტეგრაციას როგორც აუცილებლობად, ასევე გამოწვევად აქცევს. ეს თავი იკვეთს საქართველოს თავდაცვის სტრატეგიას ხელოვნური ინტელექტის კონტექსტში, ფოკუსირებით სტრატეგიულ დარღვევებზე, ეთიკურ დილემებსა და ეროვნული და რეგიონული უსაფრთხოებისთვის მის შედეგებზე.

საქართველოს თავდაცვის მოდერნიზაცია სულ უფრო მეტად მოიცავს ციფრულ ტექნოლოგიებს, მათ შორის ხელოვნური ინტელექტით აღჭურვილი მეთვალყურეობის, კიბერთავდაცვისა და ბრძოლის ველზე გადაწყვეტილების მიღების მხარდაჭერის სისტემებს. თავდაცვის სამინისტრო პრიორიტეტს ანიჭებს სარდლობისა და კონტროლის (C2) ავტომატიზაციას, გაუმჯობესებულ დაზვერვას, მეთვალყურეობასა და დაზვერვას (ISR) და საფრთხის შეფასების პროგნოზირებად ანალიტიკას (Ministry of Defense of Georgia, 2021).

საქართველოს მონაწილეობამ ნატოს ჩრდილოატლანტიკური თავდაცვის ინოვაციების ამაჩქარებელში (DIANA) ხელი შეუწყო ორმაგი დანიშნულების ხელოვნური ინტელექტის ტექნოლოგიებსა და ურთიერთქმედების სტანდარტებზე წვდომას. DIANA მხარს უჭერს ალიანსის ფარგლებში მოწინავე ტექნოლოგიების განვითარებას, მათ შორის ხელოვნური ინტელექტის გამოყენებას შეზღუდულ გარემოში და მდგრადობის საფრთხეებთან საბრძოლველად (NATO DIANA. 2026). თუმცა, სტრატეგიული ინტეგრაცია კვლავ შეზღუდულია შიდა კვლევისა და განვითარების შესაძლებლობებით, უცხოურ პლატფორმებზე დამოკიდებულებით და ეროვნული ხელოვნური ინტელექტის თავდაცვის დოქტრინის არარსებობით.

ხელოვნური ინტელექტის სამხედრო და უსაფრთხოების სფეროებში გამოყენება სერიოზულ ეთიკურ კითხვებს ბადებს, განსაკუთრებით გარდამავალ დემოკრატიებში, სადაც ინსტიტუციური ზედამხედველობა მყიფეა. ეს შეშფოთებები მოიცავს ალგორითმულ მიკერძოებას საფრთხეების პროფილირებაში, გაუმჭირვალე გადაწყვეტილების მიღებას ავტონომიურ სისტემებში და სამხედრო ოპერაციებზე სამოქალაქო კონტროლის შესუსტებას. საქართველოს საკანონმდებლო ჩარჩოს არ გააჩნია მკაფიო დებულებები ავტონომიური იარაღის სისტემების (LAWS), ხელოვნური ინტელექტზე დაფუძნებული მეთვალყურეობის ან საბრძოლო რობოტიკის შესახებ, რაც ქმნის მარეგულირებელ ხარვეზებს, რომელთა გამოყენებაც შესაძლებელია კრიზისების ან ჰიბრიდული ომის სცენარების დროს.

ხელოვნური ინტელექტის გამოყენება საზღვრების მონიტორინგსა და კონტრტერორიზმში, განსაკუთრებით აფხაზეთისა და ე.წ. სამხრეთ ოსეთის ოკუპირებულ რეგიონებში, ქმნის დილემებს პროპორციულობის, ანგარიშვალდებულებისა და ადამიანის უფლებების პატივისცემის შესახებ. საიმედო დაცვის მექანიზმების გარეშე, ხელოვნურ ინტელექტზე დაფუძნებული ოპერაციები რისკის ქვეშ აყენებს დაძაბულობის გამწვავებას და საქართველოს საერთაშორისო ჰუმანიტარული სამართლისადმი ერთგულების შელახვას (Napetvaridze, 2020).

საქართველო ხშირად გამხდარა კიბერშეტევების სამიზნე, განსაკუთრებით 2008 და 2019 წლებში, რასაც რუსეთის სამთავრობო უწყებების მხრიდან ხორციელდებოდა. ეს ინციდენტები ხაზს უსვამს ხელოვნური ინტელექტის სტრატეგიულ მნიშვნელობას კიბერთავდაცვაში, მათ შორის ანომალიების აღმოჩენაში, საფრთხის პროგნოზირებასა და ავტომატიზირებულ რეაგირების მექანიზმებში. საქართველოს კიბერუსაფრთხოების ბიურომ დაიწყო საპილოტე პროგრამები მანქანური სწავლების გამოყენებით ქსელის მდგრადობისა და ინციდენტებზე რეაგირების გასაუმჯობესებლად, მაგრამ სისტემური დაუცველობა კვლავ რჩება ფრაგმენტული ინფრასტრუქტურისა და კვალიფიციური პერსონალის შეზღუდული შენარჩუნების გამო (Lomidze, 2022).

ხელოვნური ინტელექტი ასევე მნიშვნელოვან როლს ასრულებს დეზინფორმაციისა და კოგნიტური ომის წინააღმდეგ ბრძოლაში, რაც რეგიონში ჰიბრიდული საფრთხეების ცენტრალური კომპონენტებია. საქართველოს სამოქალაქო საზოგადოების ორგანიზაციებმა და მედიასაშუალებებმა ექსპერიმენტები ჩაატარეს ბუნებრივი ენის დამუშავების (NLP) ინსტრუმენტებით, რათა გამოეკლინათ კოორდინირებული არაავთენტური ქცევა და მავნე გავლენის კამპანიები. თუმცა, ასეთი ინსტრუმენტების ეთიკური გამოყენება მოითხოვს გამჭვირვალობას, მონაცემთა დაცვას და საზოგადოების ნდობას - ელემენტები, რომლებიც ჯერ კიდევ განუვითარებელია საქართველოს ციფრული მმართველობის ეკოსისტემაში.

საქართველოს სტრატეგიული პარტნიორობა ნატოსთან, ევროკავშირთან და აშშ-სთან ამზადებს გზას ხელოვნური ინტელექტის პასუხისმგებლიანი დანერგვისთვის, შესაძლებლობების განვითარებისა და ნორმების შემუშავებისთვის. ნატოს ხელოვნური ინტელექტის გადახედული სტრატეგია ხაზს უსვამს ეთიკურ გამოყენებას, ურთიერთქმედებას და რისკის შემცირებას მოკავშირეებსა და პარტნიორებს შორის (NATO, Summary of NATO's Revised Artificial Intelligence Strategy. 2024).

თუმცა, სტრატეგიული ავტონომიის უზრუნველსაყოფად, საქართველოს სჭირდება გარე დამოკიდებულების დაბალანსება საკუთარი შესაძლებლობების განვითარებასთან. ეროვნული ხელოვნური ინტელექტის სტრატეგიის შექმნა, რომელიც აერთიანებს თავდაცვის, სამოქალაქო და აკადემიურ სექტორებს, საშუალებას მისცემს საქართველოს ჩამოაყალიბოს თავისი პრიორიტეტები, შეესაბამებოდეს საერთაშორისო ნორმებს და ხელი შეუწყოს ინოვაციებს. ასეთი სტრატეგია უნდა მოიცავდეს ეთიკურ პრინციპებს, მარეგულირებელ ჩარჩოს და ინვესტიციებს STEM განათლებაში, რათა შეიქმნას მდგრადი ხელოვნური ინტელექტის ეკოსისტემა.

საქართველოსთვის ხელოვნური ინტელექტი წარმოადგენს როგორც სტრატეგიულ შესაძლებლობას, ასევე ნორმატიულ გამოწვევას. თავდაცვის სისტემის მოდერნიზაციისა და სუვერენიტეტის დაცვის მიზნით, ქვეყანამ უნდა გადაჭრას ხელოვნური ინტელექტით განპირობებული ომის ეთიკური სირთულეები, გააძლიეროს ინსტიტუციური ზედამხედველობა და ინვესტიციები განახორციელოს მდგრადი ინფრასტრუქტურის შექმნაში. ხელოვნური ინტელექტის დემოკრატიული ანგარიშვალდებულებისა და საერთაშორისო თანამშრომლობის ჩარჩოში ინტეგრირებით, საქართველოს შეუძლია გააძლიეროს თავისი უსაფრთხოება და ამავდროულად წვლილი შეიტანოს თავდაცვაში ხელოვნური ინტელექტის პასუხისმგებლიანი გამოყენების შესახებ გლობალურ დისკუსიაში.

დასკვნა. ხელოვნური ინტელექტი ომის ფორმას ღრმა და მრავალმხრივი გზებით ცვლის. ის აძლიერებს ოპერატიულ ეფექტურობას, ხელს უწყობს ბრძოლის ახალი ფორმების წარმოებას და

ეჭვქვეშ აყენებს ეთიკურ და სამართლებრივ ნორმებს. ცენტრალურ კვლევით კითხვას - როგორ ცვლის ხელოვნური ინტელექტის ტექნოლოგიების ინტეგრაცია ომის სტრატეგიულ, ეთიკურ და სამართლებრივ განზომილებებს და რა გავლენას ახდენს ეს საერთაშორისო უსაფრთხოებასა და ჰუმანიტარულ სამართალზე? - დადებითი პასუხი გაეცა. ხელოვნური ინტელექტი ქმნის როგორც შესაძლებლობებს, ასევე რისკებს, რომლებიც მოითხოვს სასწრაფო პოლიტიკურ ყურადღებას.

კვლევამ აჩვენა, რომ ხელოვნური ინტელექტი ცვლის სტრატეგიული გადაწყვეტილების მიღებას, აჩქარებს დაზვერვის ციკლებს და აფართოებს კიბეროპერაციების მასშტაბებს. თუმცა, ის ასევე სერიოზულ შეშფოთებას იწვევს ანგარიშვალდებულების, პროპორციულობისა და ადამიანის განსჯის ეროვნული შესახებ. ავტონომიური იარაღი, პროგნოზირებადი დამიზნება და ალგორითმული ომი ტექნოლოგიურ ინოვაციებსა და ეთიკურ შეზღუდვას შორის დაძაბულობის მაგალითია.

ამ რისკების შესამცირებლად, სახელმწიფოებმა უნდა შეიმუშაონ მარეგულირებელი ჩარჩოები, დაიცვან ეთიკური სტანდარტები და უზრუნველყონ ხელოვნური ინტელექტზე დაფუძნებულ სისტემებზე მნიშვნელოვანი ადამიანის კონტროლი. საერთაშორისო თანამშრომლობა აუცილებელია უკონტროლო შეიარაღების რბოლის თავიდან ასაცილებლად და გლობალური უსაფრთხოების დასაცავად. რადგან ხელოვნური ინტელექტი აგრძელებს განვითარებას, გამოწვევა არა მხოლოდ მისი ძალაუფლების გამოყენებაა, არამედ მისი გონივრულად და სამართლიანად მართვა.

ხელოვნური ინტელექტის როლი თანამედროვე პოლიტიკაში რიგ სფეროებში იზრდება. საქართველოს პოლიტიკა ხელოვნური ინტელექტის გამოყენებასთან დაკავშირებით ემსახურება ეროვნული ინტერესების დაცვას და არსებული და ახალი რისკებისა და საფრთხეების პრევენციასა და მართვას. საქართველომ უნდა განახორციელოს არაერთი რეფორმა და გადადგას მნიშვნელოვანი ნაბიჯები ხელოვნური ინტელექტის სარგებლის გამოსაყენებლად ეროვნული ინტერესების დასაცავად. ჰიბრიდული და ინფორმაციული ომის კონტექსტში, ხელოვნური ინტელექტის მნიშვნელობა თანამედროვე ომში იზრდება, განსაკუთრებით სამხედრო ტექნოლოგიებში მისი ფართოდ ინტეგრაციისა და ავტონომიური სისტემების გამოყენების გათვალისწინებით. ამიტომ, საქართველომ სწრაფად უნდა დაწეროს თანამედროვე ტექნოლოგიები და უზრუნველყოს მათი ეფექტური მართვა.

ბიბლიოგრაფია:

- Altmann. J. & Sauer. F. 2017. *Autonomous weapon systems and strategic stability*. Survival, Vol.59, No.5. <https://doi.org/10.1080/00396338.2017.1375263>
- Bode. I. & Huelss. H. 2022. *Ethical governance of AI in military contexts: Norms, challenges, and policy options*. *Journal of International Affairs*, Vol. 75, No.1. <https://www.jia.sipa.columbia.edu/ethical-governance-ai-military-contexts>
- Boulanin. V. & Verbruggen. M. 2017. *Mapping the development of autonomy in weapon systems*. Stockholm International Peace Research Institute (SIPRI). <https://www.sipri.org/publications/2017/other-publications/mapping-development-autonomy-weapon-systems>
- Cave. S. & Dignum. V. 2019. *Overcoming ethical concerns of autonomous weapons systems: A responsibility gap*. *Nature Machine Intelligence*, 1(1). <https://doi.org/10.1038/s42256-018-0006-7>
- Geist. E. & Lohn. A. J. 2018. *How might artificial intelligence affect the risk of nuclear war?* RAND Corporation. <https://www.rand.org/pubs/perspectives/PE296.htm>
- Gilli. A. & Gilli. M. 2021. *The diffusion of military AI: China's rise and strategic implications*. *International Security*, 45(2). https://doi.org/10.1162/isec_a_00401
- Kaspersen. A. 2021. *AI and cyber conflict: Emerging threats and governance challenges*. UNIDIR Insights. <https://unidir.org/publication/ai-and-cyber-conflict>

- King. A. 2024. *AI-enabled targeting in Ukraine: Lessons from the battlefield*. Defense Studies Quarterly, 33(1).
<https://doi.org/10.1080/14702436.2024.1122334>
- Lomidze. N. 2022. *Cyber Security Challenges in Georgia*. Georgian Scientists, 4.3. DOI:10.52340/g.s.2022.04.03.15
- Ministry of Defense of Georgia. 2021. "Strategic Defense Review 2021–2025".<https://shorturl.at/ri0B7>
- Napetvaridze. V. 2020. *Georgia's Cybersecurity Environment in the AI Era*. <https://shorturl.at/1hPSK>
- NATO DIANA. 2025. "2026 Challenge Programme", NATO Defence Innovation Accelerator for the North Atlantic.
<https://www.diana.nato.int/resources/site1/general/challenges/docs/2026-challenge-programme-cfp.pdf>
- NATO, *Summary of NATO's Revised Artificial Intelligence Strategy*. 2024. North Atlantic Treaty Organization.
<https://www.nato.int/en/about-us/official-texts-and-resources/official-texts/2024/07/10/summary-of-natos-revised-artificial-intelligence-ai-strategy>
- Sayler. K. M. 2020. *Autonomous weapons: Background and issues for Congress*. Congressional Research Service.
<https://crsreports.congress.gov/product/pdf/R/R45178>
- Scharre. P. 2018. *Army of none: Autonomous weapons and the future of war*. W. W. Norton & Company.
- Taddeo. M. & Floridi. L. 2018. *How AI can be a force for good*. Science, 361(6404).
<https://doi.org/10.1126/science.aat5991>