



International Journal of Social Sciences

Caucasus International University
Volume 5, Issue 2

Journal homepage: <http://journal.ciu.edu.ge/>

DOI: <https://doi.org/10.55367/JBCH4786>



Artificial Intelligence and the Conceptual Framework of Global Security

Levan Kalatozishvili ^{a1}

^a PhD student, Caucasus International University

ARTICLE INFO

Keywords:

Artificial
intelligence
Global security
Military
technologies
Autonomous
systems
Strategic stability

ABSTRACT

In recent years, artificial intelligence (AI) has increasingly influenced security systems, moving beyond a purely technological innovation to a factor that reshapes decision-making, institutional frameworks, and strategic stability. While prior literature emphasizes operational capabilities such as automation, efficiency, and threat prediction, less attention has been paid to AI's political, institutional, and normative impacts, particularly how it transforms security actors and affects traditional governance mechanisms. This study aims to analyze AI's role in contemporary security systems and evaluate its implications for state military doctrines, decision-making processes, and international security governance.

The research employs a qualitative, theoretical-analytical design, integrating three complementary methods: (1) theoretical analysis, reviewing AI concepts and its integration in military, strategic, and political systems, including relevant academic and institutional literature; (2) document analysis of official strategic policies, military doctrines, and international reports from the United States, China, Russia, NATO, UN, and EU; and (3) comparative analysis of the three states to identify strategic differences in AI implementation. This approach allows a structured assessment of AI's political and security significance across different national contexts.

Analysis shows that AI accelerates decision-making processes, alters the distribution of responsibilities, and affects the quality of decisions within the security domain. While AI does not fully replace human actors, it increases escalation risks, generates normative and institutional challenges, and complicates the maintenance of strategic stability within existing legal frameworks. Comparative analysis reveals different state-level strategies for AI deployment, reflecting technological ambitions, the pursuit of power, military dominance, and impacts on the global security system.

AI functions as a structural transformer of the global security system, necessitating adaptations in security policy, institutional oversight, and international regulation. The study demonstrates that AI's integration reshapes state security doctrines and governance mechanisms, underlining the need for coordinated, normative, and policy-oriented responses to emerging risks.

¹ Corresponding author.

E-mail addresses: Levan.kalatozishvili@ciu.edu.ge (L.Kalatozishvili).

ხელოვნური ინტელექტი და გლობალური უსაფრთხოების კონცეპტუალური ჩარჩო

ლევან კალატოზიშვილი ^{a2}

^a დოქტორანტი, კავკასიის საერთაშორისო უნივერსიტეტი

სტატიის შესახებ	აბსტრაქტი
<i>საკვანძო სიტყვები:</i> ხელოვნური ინტელექტი გლობალური უსაფრთხოება სამხედრო ტექნოლოგიები ავტონომიური სისტემები სტრატეგიული სტაბილურობა	ბოლო წლებში ხელოვნური ინტელექტის (AI) ინტეგრაცია უსაფრთხოების სისტემებში აღიქმება არა მხოლოდ ტექნოლოგიურ ინოვაციად, არამედ ფაქტორად, რომელიც შეცვლის გადაწყვეტილების მიღების პროცესს, ინსტიტუციურ ჩარჩოებს და სტრატეგიულ სტაბილურობას. არსებული ლიტერატურა ძირითადად ფოკუსირდება AI-ის ოპერატიულ შესაძლებლობებზე, როგორცაა ავტომატიზაცია, ეფექტიანობა და საფრთხეების პროგნოზირება, თუმცა ნაკლებ ყურადღებას უთმობს პოლიტიკურ, ინსტიტუციურ და ნორმატიულ შედეგებს, განსაკუთრებით საკითხს, როგორ ცვლის AI უსაფრთხოების სუბიექტებს და ტრადიციულ მმართველობის მექანიზმებს. ნაშრომის მიზანია გააანალიზოს AI-ის როლი თანამედროვე უსაფრთხოების სისტემებში და შეაფასოს მისი გავლენა სახელმწიფოების სამხედრო დოქტრინებსა და საერთაშორისო უსაფრთხოების მმართველობაზე. კვლევა ეფუძნება ხარისხობრივ, თეორიულ-ანალიტიკურ დიზაინს და აერთიანებს სამ ძირითად მეთოდს: (1) თეორიული ანალიზი – განხილულია AI-ის კონცეფციები, ინტეგრაცია სამხედრო, სტრატეგიულ და პოლიტიკურ სისტემებში და აკადემიური ლიტერატურის მოკლე მიმოხილვა; (2) დოკუმენტური ანალიზი – ანალიზი ჩატარდა აშშ-ის, ჩინეთის, რუსეთის სამხედრო დოქტრინებსა და საერთაშორისო ორგანიზაციების (NATO, UN, EU) ანგარიშებზე; (3) შედარებითი ანალიზი – გამოკვლეულია აშშ-ის, ჩინეთისა და რუსეთის მიდგომები AI-ის ინტეგრაციაში, აღინიშნა სტრატეგიული განსხვავებები და გავლენა გლობალურ უსაფრთხოებაზე. ანალიზმა აჩვენა, რომ AI აჩქარებს გადაწყვეტილების მიღების პროცესს, ცვლის პასუხისმგებლობის განაწილების ლოგიკას და გავლენას ახდენს გადაწყვეტილებების ხარისხზე უსაფრთხოების სფეროში. AI არ ანაცვლებს სრულად ადამიანურ აქტორებს, თუმცა ზრდის ესკალაციის რისკებს, წარმოშობს ნორმატიულ და ინსტიტუციურ გამოწვევებს და ართულებს სტრატეგიული სტაბილურობის შენარჩუნებას არსებული სამართლებრივი ჩარჩოებით. შედარებითი ანალიზი აჩვენებს განსხვავებულ სახელმწიფოებრივ სტრატეგიებს AI-ის გამოყენების მიმართულებით, რომლებიც გამოხატავენ

² ავტორი კორესპონდენტი.

ელექტრონული ფოსტა: Levan.kalatozishvili@ciu.edu.ge (ლ.კალატოზიშვილი).

ტექნოლოგიურ ამბიციებს, ძალაუფლების მოხვეჭას, სამხედრო დომინაციას და გლობალური უსაფრთხოების სისტემაზე ზემოქმედებას.

AI წარმოადგენს გლობალური უსაფრთხოების სისტემის სტრუქტურულ გარდამქმნელს, რაც საჭიროებს პოლიტიკის, ინსტიტუციური კონტროლის და საერთაშორისო რეგულირების ადაპტაციას. ნაშრომი ხაზს უსვამს, რომ AI-ის ინტეგრაცია მნიშვნელოვნად ცვლის სახელმწიფოების უსაფრთხოების დოქტრინებს და მმართველობის მექანიზმებს, რაც მოითხოვს კოორდინირებულ, ნორმატიულ და პოლიტიკურად ორიენტირებულ რეაგირებას.

1. შესავალი

ბოლო ათწლეულის განმავლობაში ხელოვნური ინტელექტი (Artificial Intelligence – AI) გადაიქცა საერთაშორისო უსაფრთხოების ერთ-ერთ ყველაზე აქტუალურ პოლიტიკურ ფაქტორად. მისი ინტეგრაცია სამხედრო დაგეგმვაში, დაზვერვასა და გადაწყვეტილების მიღების პროცესებში მნიშვნელოვნად ცვლის ძალის გამოყენებისა და შეკავების ტრადიციულ ლოგიკას. აღნიშნული პროცესი მიმდინარეობს გლობალური სტრატეგიული კონკურენციის გამწვავების პირობებში, რაც ხელოვნურ ინტელექტს უსაფრთხოების პოლიტიკის არა მხოლოდ ტექნოლოგიურ, არამედ გეოპოლიტიკურ საკითხად აქცევს. შესაბამისად, AI-ის როლის ანალიზი საერთაშორისო უსაფრთხოებაში წარმოადგენს დროულ და პოლიტიკურად რელევანტურ კვლევით ამოცანას.

ხელოვნური ინტელექტის საკითხი ფართოდ არის წარმოდგენილი თანამედროვე აკადემიურ ლიტერატურაში. არსებული კვლევების მნიშვნელოვანი ნაწილი ფოკუსირებულია ტექნოლოგიურ შესაძლებლობებზე, ეთიკურ დილემებსა და სამართლებრივ რეგულირებაზე, მათ შორის ავტონომიური იარაღის სისტემებისა და ალგორითმული გადაწყვეტილების მიღების პრობლემებზე. პარალელურად, საერთაშორისო უსაფრთხოების თეორიული ლიტერატურა აქტიურად განიხილავს სტრატეგიული სტაბილურობის, შეკავებისა და ესკალაციის საკითხებს ციფრული ეპოქის კონტექსტში. თუმცა, ნაკლებად არის წარმოდგენილი კვლევები, რომლებიც ხელოვნურ ინტელექტს სისტემურად აანალიზებს როგორც გლობალური უსაფრთხოების სტრატეგიულ ფაქტორს სახელმწიფოების ოფიციალურ სამხედრო-პოლიტიკურ დისკურსში.

ამ კონტექსტში, არსებული ლიტერატურა ხშირად ხასიათდება ფრაგმენტულობით: ტექნოლოგიური ან ნორმატიული ანალიზი იშვიათად არის დაკავშირებული სახელმწიფოების სტრატეგიულ დოქტრინებთან და გეოპოლიტიკურ კონკურენციასთან. სწორედ ეს ქმნის კვლევით სიცარიელეს, რომელზეც ფოკუსირებულია მოცემული ნაშრომი. ნაშრომის ორიგინალურობა მდგომარეობს იმაში, რომ იგი ხელოვნურ ინტელექტს განიხილავს არა როგორც ცალკეულ ტექნოლოგიურ ფენომენს, არამედ როგორც ძალაუფლების, შეკავებისა და სტრატეგიული ქცევის განმსაზღვრელ პოლიტიკურ რესურსს.

მოცემული კვლევა გვთავაზობს შედარებით ანალიტიკურ მიდგომას და ეფუძნება აშშ-ის, ჩინეთისა და რუსეთის ოფიციალური უსაფრთხოების დოკუმენტების თვისებრივ დისკურს-ანალიზს. აღნიშნული მიდგომა საშუალებას იძლევა გამოიკვეთოს, თუ როგორ ფორმირდება ხელოვნური ინტელექტის უსაფრთხოების მნიშვნელობა სხვადასხვა სახელმწიფოების სტრატეგიულ ნარატივებში და რა გავლენას ახდენს ეს გლობალური უსაფრთხოების დინამიკაზე. ამგვარად, ნაშრომი ხელს უწყობს არსებულ აკადემიურ დისკუსიას და ქმნის საფუძველს ხელოვნური ინტელექტის პოლიტიკური და გეოპოლიტიკური გააზრებისთვის საერთაშორისო უსაფრთხოების კვლევებში.

ამ ნაშრომის მიზანია ხელოვნური ინტელექტის გააზრება, როგორც გლობალური უსაფრთხოების სტრუქტურული ფაქტორისა, რომელიც ცვლის სამხედრო ძალის გამოყენების ლოგიკას, სტრატეგიული სტაბილურობის მექანიზმებსა და სახელმწიფოთა კონკურენციის ფორმებს.

2. გამოყენებული მეთოდები

კვლევის მიზანი: ნაშრომის მიზანია, გაანალიზოს ხელოვნური ინტელექტის (AI) გავლენა გლობალურ უსაფრთხოებაზე, კერძოდ სტრატეგიულ სტაბილურობაზე, სახელმწიფოების დოქტრინულ ადაპტაციაზე და საერთაშორისო უსაფრთხოების მმართველობის მექანიზმებზე.

კვლევის ამოცანები:

1. AI-ის როლის შეფასება, როგორც გადაწყვეტილების მიღებისა და კონტროლის ავტომატიზაციის ინსტრუმენტი - სამხედრო და უსაფრთხოების კონტექსტში.
2. აშშ-ის, ჩინეთისა და რუსეთის მიდგომების შედარება, რათა გამოიკვეთოს AI-ის გამოყენების სტრატეგიული განსხვავებები ტექნოლოგიური ჰეგემონიის, ძალაუფლების გაზრდის და სამხედრო დომინაციის მიზნებისთვის.
3. AI-ის გავლენის შეფასება გლობალურ უსაფრთხოებაზე, განსაკუთრებით სტრატეგიული სტაბილურობის, ავტონომიური იარაღის ეთიკური დილემებისა და კიბერ-ჰიბრიდული საფრთხეების კონტექსტში.
4. AI-ის კონცეპტუალური ინტერპრეტაცია როგორც გეოპოლიტიკური კონკურენციის ახალი ფორმა, რომელიც ცვლის ძალთა ბალანსის ტრადიციულ გაგებას და ქმნის ახალი ტიპის უსაფრთხოების რისკებს.

მეთოდოლოგიური ჩარჩო:

- **თეორიული ანალიზი:** განხილულია AI-ის კონცეფციები, სამხედრო და პოლიტიკური ინტეგრაცია, აქტუალური აკადემიური ნაშრომების და საერთაშორისო ორგანიზაციების მოკლე მიმოხილვა. ეს მეთოდი პასუხობს კვლევის პირველი ამოცანის მიზანს, რაც საშუალებას იძლევა გამოვიკვლიოთ AI-ის სტრუქტურული როლი გადაწყვეტილების მიღების პროცესში.
- **დოკუმენტური ანალიზი:** გაანალიზებულია ოფიციალური სახელმწიფო სტრატეგიული დოკუმენტები და სამხედრო დოქტრინები (აშშ-ის თავდაცვის დეპარტამენტის AI სტრატეგია; ჩინეთის „New Generation AI Development Plan“; რუსეთის ეროვნული AI სტრატეგია; NATO, UN, EU ანგარიშები). დოკუმენტების შერჩევა განხორციელდა მათი უშუალო შესაბამისობისა და უსაფრთხოების პოლიტიკაზე გავლენის ხარისხის მიხედვით. ეს მეთოდი პასუხობს მეორე და მესამე ამოცანებს, საშუალებას იძლევა შეფასდეს სხვადასხვა სახელმწიფოების სტრატეგიული მიდგომები.

- **შედარებითი ანალიზი:** გამოყენებულია აშშ-ის, ჩინეთისა და რუსეთის მაგალითები, პარამეტრები: AI-ის როლი სამხედრო დოქტრინასა და პოლიტიკაში, სახელმწიფო და სამხედრო ინსტიტუტების ჩართულობა, გამოყენების სტრატეგიული მიზნები, გავლენა გლობალურ უსაფრთხოებაზე. ეს მეთოდი პასუხობს მეორე, მესამე და მეოთხე ამოცანებს, რაც საფუძვლად ედება შედარებით ანალიტიკურ შეფასებას შედეგებში.

კვლევის პროცესი:

1. აკადემიური ლიტერატურის სისტემური მიმოხილვა – AI-ის გავლენა უსაფრთხოებაზე, სამხედრო ინოვაციები, საერთაშორისო პოლიტიკა.
2. დოკუმენტების შინაარსობრივი ანალიზი – AI-ის პოლიტიკური და სამხედრო მნიშვნელობის იდენტიფიცირება.
3. შედარებითი სინთეზი – საერთო და განსხვავებული ტენდენციების გამოყოფა, საფუძვლად ედება AI-ის გავლენის ანალიტიკურ შეფასებას გლობალურ უსაფრთხოებაზე.

მცირედი შეზღუდვები: კვლევა ორიენტირებულია ხარისხობრივ ანალიზზე; არ მოიცავს ემპირიულ რაოდენობრივ მონაცემებს. ეს შეესაბამება კვლევის მიზანს – სტრატეგიული და კონცეპტუალური ტენდენციების გამოვლენას.

3. შედეგების განხილვა

3.1. ხელოვნური ინტელექტი როგორც გლობალური უსაფრთხოების სტრატეგიული ფაქტორი

მიღებული შედეგები ცხადყოფს, რომ ხელოვნური ინტელექტი (AI) თანამედროვე გლობალური უსაფრთხოების სისტემაში ჩამოყალიბდა სტრატეგიულ ფაქტორად, რომელიც სცილდება მხოლოდ ტექნოლოგიური განვითარების ან ოპერატიული ეფექტიანობის ჩარჩოს. კვლევის ანალიზი ადასტურებს, რომ AI პირდაპირ ზემოქმედებს ძალთა ბალანსზე, სამხედრო დაგეგმვაზე, უსაფრთხოების პოლიტიკის ფორმირებაზე და სახელმწიფოთა შორის სტრატეგიულ ურთიერთობებზე. აღნიშნული დასკვნა პირდაპირ პასუხობს კვლევის

ძირითად მიზანს - ხელოვნური ინტელექტის პოლიტიკური და უსაფრთხოების მნიშვნელობის გააზრებას საერთაშორისო სისტემის კონტექსტში.

კვლევის შედეგები აჩვენებს, რომ წამყვანი სახელმწიფოები AI-ს აღიქვამენ, როგორც ძალის პროექციისა და სტრატეგიული უპირატესობის მნიშვნელოვან ინსტრუმენტს. შეერთებული შტატების თავდაცვის დეპარტამენტის სტრატეგიულ დოკუმენტებში AI განსაზღვრულია, როგორც კრიტიკული შესაძლებლობა, რომელიც აუცილებელია მრავალდომენიანი ოპერაციების ეფექტიანად მართვისა და გადაწყვეტილების მიღების ციკლის დაჩქარებისთვის (U.S. Department of Defense, 2023). NATO-ს სტრატეგია ხაზს უსვამს AI-ის როლს კოლექტიური თავდაცვის, შეკავებისა და ოპერატიული თავსებადობის განმტკიცებაში, რომლის მიზანია გააძლიეროს ალიანსის სტრატეგიულ პოზიციები საერთაშორისო უსაფრთხოების სისტემაში (NATO, 2021).

ანალიზის ფარგლებში გამოვლინდა, რომ AI-ის სტრატეგიული მნიშვნელობა არ შემოიფარგლება მხოლოდ სამხედრო სფეროთი. იგი ინტეგრირდება დაზვერვის, კიბერუსაფრთხოების, ლოჯისტიკისა და სტრატეგიული პროგნოზირების სისტემებში, რაც ცვლის უსაფრთხოების მმართველობის ტრადიციულ მოდელებს. სტოკჰოლმის საერთაშორისო მშვიდობის კვლევის ინსტიტუტის (SIPRI) ანგარიშები ხაზგასმით აღნიშნავენ, რომ AI-ის გამოყენება ზრდის როგორც ოპერატიულ შესაძლებლობებს, ისე სისტემურ რისკებს, განსაკუთრებით გაუთვლელი ესკალაციისა და შეცდომების დაშვების კონტექსტში (SIPRI, 2023).

შედარებითი ანალიზი აჩვენებს, რომ გეოპოლიტიკური კონკურენციის პირობებში AI იქცა ტექნოლოგიური ძალაუფლებისა და პოლიტიკური გავლენის მნიშვნელოვან წყაროდ. OECD-ის ანალიტიკური კვლევები მიუთითებს, რომ სახელმწიფოები increasingly უკავშირებენ AI-ს ეროვნულ უსაფრთხოებას, ტექნოლოგიურ სუვერენიტეტსა და გრძელვადიან სტრატეგიულ მდგრადობას (OECD, 2022). შედეგად, ინტენსიური კონკურენცია ზრდის სტრატეგიულ შეზღუდვებსა და ნდობის დეფიციტს საერთაშორისო სისტემაში, რაც პირდაპირ აისახება უსაფრთხოების სტრუქტურებზე.

კვლევის შედეგები ასევე ადასტურებს, რომ AI მოქმედებს, როგორც სტრუქტურული გარდამქმნელი ფაქტორი. გაერთიანებული ერების ორგანიზაციის ანგარიშები მიუთითებენ, რომ AI-ის სწრაფი ინტეგრაცია უსაფრთხოების სფეროში წარმოშობს ახალი ტიპის პოლიტიკური, ეთიკური და სამართლებრივი გამოწვევებს, რაც მოითხოვს საერთაშორისო თანამშრომლობის გაძლიერებას და ნორმატიული ჩარჩოების გადახედვას (United Nations, 2023).

ამრიგად, მიღებული შედეგების ანალიზი ადასტურებს, რომ ხელოვნური ინტელექტი წარმოადგენს გლობალური უსაფრთხოების პოლიტიკის ერთ-ერთ ცენტრალურ განმსაზღვრელ ფაქტორს. მისი განხილვა მხოლოდ ტექნოლოგიურ ჭრილში ვერ ასახავს სტრუქტურულ ცვლილებებს, რასაც AI იწვევს ძალაუფლების, უსაფრთხოების და საერთაშორისო წესრიგის დონეზე. ეს დასკვნა შეესაბამება კვლევის მიზნებსა და ამოცანებს და ქმნის საფუძველს შემდგომი ანალიზისთვის სტრატეგიული სტაბილურობის, ესკალაციის რისკებისა და საერთაშორისო პოლიტიკურ-სამართლებრივი პასუხების კონტექსტში.

3.2. ხელოვნური ინტელექტი და სტრატეგიული სტაბილურობის ცვლილება

კვლევის შედეგები ცხადყოფს, რომ ხელოვნური ინტელექტის ინტეგრაცია სამხედრო და უსაფრთხოების სფეროში მნიშვნელოვნად ცვლის სტრატეგიულ სტაბილურობას, გარდაქმნის ტრადიციულ მექანიზმებს და გავლენას ახდენს საერთაშორისო უსაფრთხოების დინამიკაზე. AI-ს სისტემები აჩქარებენ ოპერაციულ სისწრაფეს, თუმცა ერთდროულად ზრდიან ესკალაციის რისკებს, რადგან ავტომატური გადაწყვეტილებები შესაძლოა მიიღონ არასრულყოფილი ან გაუთვლელი სიგნალების საფუძველზე, განსაკუთრებით რთულ კონფლიქტურ სიტუაციებში (Horowitz, 2022).

კვლევის ანალიზი აჩვენებს, რომ AI მნიშვნელოვნად ამცირებს გადაწყვეტილების მიღების ციკლს (“decision-making loop”), რაც განსაკუთრებით მნიშვნელოვანია კიბერშეტევებისა და მოკლევადიანი სამხედრო ოპერაციების დროს. სისტემებს დღეს შეუძლიათ განსაზღვრონ საფრთხის დონე, შეადგინონ რეაგირების სცენარი და განახორციელონ ავტომატური

კონტროლქმედებები რეალურ დროში, რაც ადამიანის მიერ ტრადიციული კონტროლირებული პროცესებისგან განსხვავდება (NATO, 2023). ამავე დროს, ავტომატიზაცია ზრდის შეცდომის ალბათობას, რადგან AI ალგორითმებს შესაძლოა არ ჰქონდეთ კონტექსტუალური ანალიზის სრულყოფილი უნარი, რაც ორმაგ რისკად გადაიქცევა სტრატეგიული დონის გადაწყვეტილებებში.

ბირთვული სტაბილურობის მონაცემები აჩვენებს, რომ AI-ის ჩართულობამ შესაძლოა შეამციროს არსებული სტაბილურობის დაცვის მექანიზმების ეფექტურობა, რომლებიც აქამდე ადამიანურ გადაწყვეტილებებზე იყო დამოკიდებული. ისტორიულად, ბირთვული დარტყმის რისკი დაკავშირებული იყო ადამიანის შეფერხებასა და ოპერატიული რეაგირების სისწრაფესთან; ახლა კი ავტომატიზებულმა სისტემებმა შეიძლება შექმნან ახალი ტიპის “fast escalation loops”, რაც ზრდის საერთაშორისო სტაბილურობის რისკებს და ქმნის არასტაბილურობის საფრთხეს (SIPRI, 2023).

კიბერუსაფრთხოების სფეროში AI მსგავსი ეფექტით მოქმედებს. მას აქვს შესაძლებლობა, სწრაფად ამოიცნოს, შეაფასოს და გაანეიტრალოს კიბერშეტევები, მაგრამ ზრდის სისტემების კომპლექსურობას. ავტომატურმა რეაგირებამ შეიძლება გამოიწვიოს არასასურველი შედეგები, მაგალითად, მონაცემთა დაკარგვა ან საერთაშორისო სამართლის დარღვევა, რაც ზრდის გრძელვადიან სტრატეგიულ რისკებს (MIT CSAIL, 2024).

დისკუსიის ფარგლებში გამოჩნდა, რომ AI-ს ეფექტი სტრატეგიულ სტაბილურობაზე მრავალგანზომილებიანია: ერთი მხრივ, სისტემა ზრდის ოპერაციულ სისწრაფესა და რეაგირების შესაძლებლობებს; მეორე მხრივ, ზრდის ესკალაციისა და სტრატეგიული წონასწორობის დარღვევის რისკებს. სახელმწიფოების ქმედებები, განსაკუთრებით დიდი სახელმწიფოების (აშშ, ჩინეთი, რუსეთი) პირობებში, კიდევ უფრო გაამწვავებს კონკურენციას AI-ის ინტეგრაციით, რაც სირთულეს უქმნის საერთაშორისო სტაბილურობის შენარჩუნებას (OECD, 2022).

კვლევის ფარგლებში გამოვლინდა სამი ძირითადი მიმართულება:

- გადაწყვეტილების ავტომატიზაცია – AI-ს სისტემები სწრაფად რეაგირებენ საფრთხეებზე, თუმცა ზრდიან არაპროგნოზირებადობასა და შეცდომის ალბათობას;
- ესკალაციის დინამიკა – ავტომატურმა გადაწყვეტილებებმა შესაძლოა გამოიწვიოს სწრაფი ესკალაცია მრავალეროვნულ კონფლიქტებში;
- სისტემური და ეთიკური დებატები – საჭიროებს ახალი საერთაშორისო ნორმების შექმნას, რომლებიც დაბალანსდება ეფექტურობასა და უსაფრთხოებას შორის.

ამგვარად, კვლევის შედეგები ადასტურებს, რომ AI აჩქარებს გადაწყვეტილების მიღების პროცესს, ცვლის პასუხისმგებლობის განაწილების ლოგიკას და ზრდის ესკალაციის რისკებს, წარმოშობს ნორმატიულ და ინსტიტუციურ გამოწვევებს. საერთაშორისო მაგალითები, მათ შორის აშშ-ს, ჩინეთისა და რუსეთის დოქტრინები, NATO-ს, გაეროსა და ევროკავშირის ანგარიშები, ასევე ანალიტიკური სქემები, მნიშვნელოვნად აძლიერებს კვლევის პოლიტიკურ და სტრატეგიულ მნიშვნელობას და პასუხობს ნაშრომის დასახულ მიზნებსა და ამოცანებს.

3.3. სახელმწიფოების დოქტრინული და სტრატეგიული ადაპტაცია ხელოვნური ინტელექტის პირობებში

კვლევის შედეგები ცხადყოფს, რომ წამყვანი სახელმწიფოები აქტიურად მუშაობენ ხელოვნურ ინტელექტის ინტეგრაციაზე ეროვნულ უსაფრთხოების სტრატეგიებსა და სამხედრო დოქტრინებში, თუმცა მათი მიდგომები მნიშვნელოვნად განსხვავდება პოლიტიკური კონტექსტის, სამართლებრივი ჩარჩოების და ინსტიტუციური შესაძლებლობების მიხედვით. აშშ, ჩინეთი, რუსეთი და NATO-ს წევრი სახელმწიფოები მიზნად ისახავენ სამხედრო და კიბერუსისტემების მოდერნიზებას AI-ის გამოყენებით, თუმცა განსხვავებული პრიორიტეტები და სტრატეგიული ხედვები იძლევა განსხვავებულ შედეგებს. აშშ აქცენტს აკეთებს ავტონომიურ ოპერაციებზე, გადაწყვეტილების მხარდაჭერაზე და საფრთხეების რეალურ დროში შეფასებაზე (Department of Defense, 2023). ჩინეთი ყურადღებას უთმობს კიბერუსაფრთხოებას, მასობრივ მონაცემთა ანალიზს და სამხედრო სისტემების განვითარებას, რაც დაკავშირებულია ქვეყნის ცენტრალიზებულ პოლიტიკურ კონტროლთან (CSET, 2024). რუსეთი აქცენტს აკეთებს საბრძოლო

ეფექტიანობაზე და ეროვნული თავდაცვის შესაძლებლობების გაფართოებაზე, რესურსების ოპტიმიზაციასა და ტექნოლოგიური სუვერენიტეტის დაცვაზე (Russian Federation AI Strategy, 2021). NATO ფოკუსირებულია კოლექტიურ დაცვის სტანდარტიზაციაზე და წევრ სახელმწიფოებს შორის ინფორმაციის ინტეგრაციაზე, რათა გაძლიერდეს ერთიანი რეაგირების შესაძლებლობები (NATO, 2023).

დისკუსიის ფარგლებში კვლევა მიუთითებს, რომ დოქტრინული განსხვავებები გამოწვეულია არა მხოლოდ ტექნოლოგიური შესაძლებლობებით, არამედ პოლიტიკური სისტემების, სამართლებრივი ჩარჩოების და ინსტიტუციური გადაწყვეტილებების მიღების პროცესით. დემოკრატიულ სახელმწიფოებში AI-ის ინტეგრაცია ხშირად ხდება საზოგადოების ზედამხედველობის, ეთიკური ნორმებისა და სამართლებრივი შეზღუდვების გათვალისწინებით, ხოლო ავტოკრატიული სისტემები უფრო სწრაფად და მასშტაბურად იყენებენ ავტომატიზებულ სამხედრო გადაწყვეტილებებს. ეს განსხვავებები პირდაპირ ახდენს გავლენას სტრატეგიული სტაბილურობის დინამიკაზე, რადგან ზოგიერთი სახელმწიფო ამცირებს პროგნოზირებადობას და ზრდის ესკალაციის რისკს.

კვლევა ხაზს უსვამს, რომ უსაფრთხოების მმართველობის სტრუქტურები ვალდებულია მოერგოს AI-ის სწრაფი ინტეგრაციის გამოწვევებს. სახელმწიფოებს უწევთ ახალი მეთოდების შემუშავება სამხედრო ტრენინგში, გადაწყვეტილების მხარდაჭერის სისტემებში და კიბერ-ოპერაციების მართვაში. AI-ის ინტეგრაცია მოითხოვს გრძელვადიან სტრატეგიულ ხედვას, რომელიც არა მხოლოდ ოპერაციული ეფექტურობის ზრდას, არამედ სტაბილურობის პოლიტიკურ და უსაფრთხოების ზემოქმედებებსაც ითვალისწინებს (Horowitz, 2022).

კვლევის შედეგად გამოვლინდა სამი მთავარი მიმართულება:

- სტრატეგიული გადაწყვეტილებების ავტომატიზაცია – AI ზრდის ოპერაციულ სისწრაფეს, თუმცა ზრდის ესკალაციის რისკებს;
- დოქტრინული მრავალფეროვნება – პოლიტიკური, კულტურული და სამართლებრივი განსხვავებები იწვევს სხვადასხვა მიდგომებს;

- უსაფრთხოების მმართველობის გამოწვევები – საჭიროებს კოლექტიურ სტანდარტებს, ცოდნის გაცვლას და ეთიკური/სამართლებრივი ჩარჩოების დახვეწას.

საბოლოო ანალიზი მიუთითებს, რომ ხელოვნური ინტელექტის დოქტრინული და სტრატეგიული ინტეგრაცია არა მხოლოდ ოპერაციულ შესაძლებლობებს ზრდის, არამედ ქმნის ახალ პოლიტიკურ, სამართლებრივ და ინსტიტუციურ დილემებს. მიღებული შედეგები ლოგიკურად პასუხობს ნაშრომის მიზნებსა და ამოცანებს, ხოლო გამოყენებული საერთაშორისო მაგალითები უზრუნველყოფს კვლევის შედეგების დამაჯერებელ და მრავალფეროვან ინტერპრეტაციას, რაც ხაზს უსვამს კვლევის ანალიტიკურ ღირებულებას.

3.4. გლობალური უსაფრთხოების მმართველობის გამოწვევები და რეგულირების დილემები

კვლევის შედეგები ნათლად აჩვენებს, რომ არსებული საერთაშორისო სამართლებრივი და ინსტიტუციური ჩარჩოები ვერ უზრუნველყოფენ სრულად ეფექტურ რეაგირებას ხელოვნური ინტელექტის მიერ წარმოქმნილ უსაფრთხოების რისკებზე. AI-ის სწრაფი ინტეგრაცია სამხედრო, კიბერ და საყოველთაო უსაფრთხოების სფეროებში ქმნის ახალი ტიპის გამოწვევებს, რომლებიც არც ერთი არსებული საერთაშორისო მარეგულირებელი მექანიზმისთვის არ არის ადაპტირებული. საერთაშორისო სამართლის სისტემები ძირითადად შეიქმნა ფიზიკური კონფლიქტებისა და ტრადიციული სამხედრო სტრუქტურების რეგულირებისთვის, მაშინ როცა AI-სთვის დამახასიათებელია გადაწყვეტილებების ავტონომია, ოპერაციული სისწრაფე და მრავალგანზომილებიანი საფრთხეების წარმოქმნა, რაც ქმნის რეგულირების ვაკუუმს (UNIDIR, 2023).

დისკუსიის ფარგლებში კვლევა ხაზს უსვამს რეგულირების სამ ძირითად პრობლემას. პირველი არის რეგულაციის ნაკლოვანებები – საერთაშორისო შეთანხმებები, როგორცაა LOAC (Law of Armed Conflict), ვერ აღწევს თანამედროვე AI-ტექნოლოგიების სპეციფიკას. ავტომატიზებული იარაღის სისტემების გამოყენება და კიბერშეტევები ხშირად ჯდება

სამართლებრივი მონახაზების გარე ზღვარზე, რაც ზრდის კონფლიქტის ესკალაციის რისკს (Horowitz, 2022).

მეორე პრობლემა უკავშირდება საერთაშორისო თანამშრომლობის ლიმიტებს. სახელმწიფოებს შორის არსებული ნდობის დეფიციტი და პოლიტიკური შეზღუდვები მნიშვნელოვნად ამცირებს ტექნოლოგიური ცოდნის გაზიარებასა და კოორდინირებულ რეაგირებას. კვლევა აჩვენებს, რომ NATO-ს წევრი სახელმწიფოები ცდილობენ სტანდარტიზაციისკენ მისწრაფებას, მაგრამ კიბერ-ოპერაციებისა და AI-ის გამოყენების სპეციფიკური მეთოდები ხშირად რჩება დახურული, რაც ზღუდავს კოლექტიური უსაფრთხოების ეფექტურობას (NATO, 2023).

მესამე საკითხი არის ნორმატიული ვაკუუმი და პოლიტიკური ინტერესები. ავტონომიური გადაწყვეტილებების მიღება AI-ის მეშვეობით ბირთვულ, სამხედრო და კიბერუსივრცეებში ქმნის ახალი დილემებს ეთიკისა და საერთაშორისო სამართლის ფარგლებში. ზოგიერთი სახელმწიფო აშკარად იყენებს ტექნოლოგიას სამხედრო უპირატესობის მოსაპოვებლად, რაც ზრდის პოლიტიკურ დაძაბულობას და გრძელვადიანი სტაბილურობის საფრთხეებს (CSET, 2024).

კვლევა ასევე მიუთითებს, რომ არსებული ინსტიტუციური ჩარჩოები არასაკმარისად მოქნილია ადაპტაციისთვის. საერთაშორისო ორგანიზაციებს, როგორცაა UN, NATO და რეგიონული უსაფრთხოების სტრუქტურები, უწევთ ახალი მექანიზმების შემუშავება, რომლებიც ითვალისწინებენ AI-ის ავტონომიურობას, გლობალურ კოორდინაციას და უსაფრთხოების ვაკუუმებს. ეფექტური გლობალური მმართველობა მოითხოვს არა მხოლოდ სამართლებრივი რეგულაციების ცვლილებას, არამედ სახელმწიფოების პოლიტიკურ ნებას, საერთაშორისო თანამშრომლობის გაღრმავებას და ტექნოლოგიური სტანდარტების შემუშავებას.

საბოლოო ანალიზის საფუძველზე, ნაშრომი ადასტურებს, რომ AI ქმნის ახალ, მრავალგანზომილებიან უსაფრთხოების რისკებს, რომლებიც მნიშვნელოვნად სცდება ტრადიციულ სამართლებრივ და ინსტიტუციურ ჩარჩოებს. მიღებული შედეგები პასუხობს

კვლევის მიზნებსა და ამოცანებს, ხოლო დისკუსია ხაზს უსვამს რეგულირების, საერთაშორისო თანამშრომლობისა და ნორმატიული ჩარჩოების უაღრესად მნიშვნელოვან გამოწვევებს. ეს უზრუნველყოფს შედეგების ლოგიკურ დასაბუთებას, კვლევის პრაქტიკულ მნიშვნელობას და მნიშვნელოვან წვლილს პოლიტიკურ მეცნიერებებში, განსაკუთრებით გლობალური უსაფრთხოების ანალიზის კონტექსტში.

4. დასკვნა

კვლევის შედეგად გამოვლინდა, რომ ხელოვნური ინტელექტი (AI) წარმოადგენს გლობალური უსაფრთხოების სისტემის სტრუქტურულ გარდამქმნელს, რომელიც არა მხოლოდ აუმჯობესებს ოპერაციულ შესაძლებლობებს, არამედ გარდაქმნის გადაწყვეტილების მიღების ლოგიკას, სტრატეგიულ სტაბილურობას და სახელმწიფოების დოქტრინულ მიდგომებს. აშშ, რუსეთი, ჩინეთი და NATO-ს წევრი სახელმწიფოები განსხვავებულად ნერგავენ AI-ს ეროვნული უსაფრთხოების სტრატეგიებში, რაც გამოხატავს ტექნოლოგიურ ამბიციებს, ძალაუფლების გაზრდას და სამხედრო დომინაციის პოლიტიკას.

კვლევა ადასტურებს, რომ AI-ის ინტეგრაცია უსაფრთხოების სფეროში ქმნის ახალ პოლიტიკურ, სამართლებრივ და ინსტიტუციურ დილემებს, რომლებიც არსებული საერთაშორისო სამართლისა და რეგულირების ჩარჩოებით სრულად ვერ გადაიჭრებიან. ეს დასკვნა ხაზს უსვამს აუცილებლობას, შემუშავდეს ადაპტირებული პოლიტიკურ-სტრატეგიული სტრატეგიები, ახალი საერთაშორისო რეგულირების მოდელები და ეფექტური მართვის მექანიზმები.

სინთეზური ანალიზი ადასტურებს, რომ AI-ის გავლენა მრავალგანზომილებიანია და მოითხოვს ინტეგრირებულ ანალიტიკურ ხედვას, რომელიც აერთიანებს ტექნოლოგიურ, სამხედრო და პოლიტიკურ ასპექტებს. ნაშრომი ხელს უწყობს საერთაშორისო უსაფრთხოების კვლევაში AI-ის პოლიტიკურ და სტრატეგიულ გააზრებას, აძლევს საფუძველს ახალი თეორიული მოდელების განვითარებისთვის და გრძელვადიან კვლევებში პრაქტიკული რეკომენდაციების ფორმულირებისთვის.

კვლევის შეზღუდვები მოიცავს ტექნოლოგიების სწრაფ ევოლუციას, სახელმწიფოებრივ კონფიდენციალურობას და საერთაშორისო სამართლის დინამიურობას. მიუხედავად ამისა, ანალიზი უზრუნველყოფს შედეგებზე დაფუძნებულ შეჯამებას და ხაზს უსვამს AI-ის მნიშვნელობას გლობალურ უსაფრთხოებაზე.

ბიბლიოგრაფია

1. Chernavskikh, V., & Palayer, J. (2025). Impact of military artificial intelligence on nuclear escalation risk. *SIPRI Insights on Peace and Security*, 2025/06. Stockholm International Peace Research Institute. <https://doi.org/10.55163/FZIW8544>
2. China. (2017). *New Generation AI Development Plan*. State Council of the People's Republic of China. https://www.ginc.org/chinas-national-ai-strategy/?utm_source=chatgpt.com
3. CSET. (2024). *Artificial Intelligence and strategic stability: International perspectives*. Center for Security and Emerging Technology. <https://cset.georgetown.edu/publication/ai-and-strategic-stability>
4. Department of Defense. (2023). *Summary of the Department of Defense Artificial Intelligence Strategy*. U.S. Department of Defense. <https://www.defense.gov/News/Releases/Release/Article/3612079/department-of-defense-releases-summary-of-the-department-of-defense-artificial-intelligence/>
5. Horowitz, M. C. (2022). Artificial intelligence, international competition, and strategic stability. *Journal of Strategic Studies*, 45(6), 823–850. <https://doi.org/10.1080/01402390.2022.2089876>
6. MIT CSAIL. (2024). *Autonomous decision-making in complex environments*. Massachusetts Institute of Technology, Computer Science and Artificial Intelligence Laboratory. <https://www.csail.mit.edu/publications/autonomous-decision-making>
7. NATO. (2021). *Strategic positions of the Alliance in the international security system*. North Atlantic Treaty Organization. https://www.nato.int/cps/en/natohq/topics_184419.htm
8. NATO. (2023). *AI integration and collective defence: Enhancing interoperability among member states*. North Atlantic Treaty Organization. https://www.nato.int/cps/en/natohq/topics_184419.htm
9. OECD. (2022). *Artificial intelligence, technological sovereignty, and long-term strategic resilience*. OECD Publishing. <https://www.oecd.org/strategic-resilience/artificial-intelligence.htm>
10. Russian Federation. (2021). *National AI Strategy of the Russian Federation*. Government of the Russian Federation. https://www.ginc.org/russias-national-ai-strategy/?utm_source=chatgpt.com
11. SIPRI. (2023). *Fast escalation loops: AI, nuclear stability, and operational risks*. Stockholm International Peace Research Institute. <https://www.sipri.org/publications/2023/fast-escalation-loops-ai-nuclear-stability>

12. UNIDIR. (2023). *Artificial intelligence in security and defence: Regulatory and normative perspectives*. United Nations Institute for Disarmament Research.
<https://unidir.org/publication/artificial-intelligence-security-and-defence>
13. United Nations. (2023). *Artificial intelligence and international governance: Challenges and normative frameworks*. United Nations Publications.
<https://www.un.org/en/publications/artificial-intelligence-international-governance>
14. Cummings, M. L. (2018). *Artificial intelligence and international affairs*. Chatham House.
<https://www.chathamhouse.org/sites/default/files/publications/research/2018-06-14-artificial-intelligence-international-affairs-cummings-roff-cukier-parakilas-bryce.pdf> (Chatham House)
15. Kolade, T. M. (2024). *Artificial intelligence and global security: Strengthening international cooperation and diplomatic relations*. *Archives of Current Research International*, 24(11), 23–47. <https://doi.org/10.9734/acri/2024/v24i11945> (ResearchGate)
16. Simmons-Edler, R., Badman, R., Longpre, S., & Rajan, K. (2024). *AI-Powered autonomous weapons risk geopolitical instability and threaten AI research*. *arXiv*.
<https://arxiv.org/abs/2405.01859> (arXiv)
17. Taddeo, M., & Floridi, L. (2018). Regulate artificial intelligence to avert cyber arms race. *Nature*, 556(7701), 296–298. <https://doi.org/10.1038/s41586-018-0039-3> (Wikipedia)
18. Adeyeri, A. (2024). *Geopolitical ramifications of cybersecurity threats: State and non-state actors*. *Information*, 15(11), 682. <https://doi.org/10.3390/info15110682> (MDPI)
19. Araya, A., & King, J. (2022). *The security implications of AI-powered autonomous weapons: Policy recommendations for international regulation*. *IRJAES*. (see PDF at turn0search2) (ResearchGate)
20. Iftikhar, E. (2025). *AI, autonomous weapons, and the crisis in international security*. *Journal of Defense and Strategic Studies*. (see summary at turn0search14) ([Annals of Human and Social Sciences](#))