

Süni İntellektin Hərbi Sahədə Tətbiqi: Qlobal Təhlükəsizlik və Kibermüdafiə Strategiyaları

Nuran Mahmudov¹, Mehrac Hüseynov²

Milli Müdafiə Universiteti, Bakı, Azərbaycan

¹mahmudovnuran@outlook.com, ²mehrac77@gmail.com

Xülasə— Bu tədqiqat materialı süni intellektin hərbişdirilməsi prosesinin və qlobal təhlükəsizlik sistemlərinə transformasiyaedici təsirinin çoxölçülü təhlilini təqdim edən hərtərəfli hesabatdır. Tədqiqatın əsasını Hərbi Qərar Qəbuletmə Prosesinin və OODA ((Observe, Orient, Decide, Act) dövrünün optimallaşdırılması təşkil edir. Rəqib maşın öyrənməsi vasitəsilə həyata keçirilən təhdidlər, izah edilə bilən Sİ arxitekturaları, avtomatlaşdırma qərəzliliyi və beynəlxalq humanitar hüquq normaları çərçivəsində "mənalı insan nəzarəti" ((Meaningful Human Control)) konsepti dəyərləndirilmişdir. Sonda Azərbaycan Respublikasının kibertəhlükəsizlik strategiyası, Milli Təhlükəsizlik Əməliyyatları Mərkəzinin rolu və ölkənin rəqəmsal suverenliyinin təmin edilməsi məsələləri elmi və praktiki aspektlərdən təhlil edilmişdir.

Açar sözlər— süni intellekt; hərbişdirmə; kibertəhlükəsizlik; avtonom silah sistemləri; maşın öyrənməsi.

I. GİRİŞ

Müasir qlobal təhlükəsizlik memarlığı bəşəriyyət tarixində görünməmiş bir texnoloji transformasiyadan keçir. Rəqəmsal inqilabın ən mühüm innovasiyalarından biri olan süni intellekt (Sİ) artıq sadəcə mülki və kommersiya məqsədli texnologiya deyil, misilsiz miqyasda, mürəkkəblikdə və təsir gücündə olan bir hərbi silaha çevrilmişdir. Bu hərbişdirmə prosesi ənənəvi müharibə qaydalarını tamamilə yenidən yazır və dövlətlərin təhlükəsizlik doktrinalarını kökündən dəyişdirir. Təbqasına qərar verə bilən avtonom platformalardan tutmuş, qlobal kiberməkanda saniyənin mində biri qədər qısa müddətdə həyata keçirilən kütləvi hücum kampaniyalarına qədər geniş bir spektri əhatə edən Sİ, artıq beynəlxalq münasibətlər sistemində güc balansının əsas təyinedici faktorudur.

Süni intellektin hərbi sahədəki tətbiqlərinin fundamental əsasını "Sİ triadası" adlandırılan üç kritik komponentin – alqoritmlərin, məlumat bazalarının və hesablama gücünün misli görünməmiş sinerjisi təşkil edir. Bu üç komponentin bir araya gəlməsi "alqoritmik müharibə" erasının başlanğıcını qoyaraq, Şimali Atlantika Müqaviləsi Təşkilatı (NATO) da daxil olmaqla, bir çox hərbi alyansların və milli dövlətlərin öz müdafiə və hücum strategiyalarını yenidən formalaşdırmasını zəruri etmişdir. Alqoritmik müharibənin əsas hədəfi düşmən üzərində kinetik üstünlük əldə etməkdən daha çox, informasiya və qərar qəbuletmə sürəti üzərində mütləq üstünlük qurmaqdır. Bu yanaşma sübut edir ki, gələcəyin müharibələrində qalib gələn

tərəf daha çox əsgərə və ya ənənəvi silaha sahib olan tərəf deyil, məlumatı daha sürətli və daha dəqiq analiz edərək hərəkətə keçə bilən tərəf olacaqdır.

Problemin aktuallığı ondan ibarətdir ki, süni intellekt sistemləri hərbi əməliyyatların və kiber hücumların sürətini insan idrakının hüdudlarından kənara çıxarır. Məlumdur ki, hərbi idarəetmənin mərkəzində duran qərarların qəbul edilməsi və onların icrası arasındakı zaman fərqi (latency) Sİ vasitəsilə minimuma, hətta mikrosaniyələrə endirilir. Lakin bu misilsiz sürət və miqyas avtomatlaşdırma qərəzliliyi (automation bias), məlumat zəhərlənməsi (data poisoning), "qara qutu" (black box) qeyri-şəffaflığı və hesabatlılığın itməsi kimi fundamental riskləri də özündə ehtiva edir. Xüsusilə, Avtonom Silah Sistemlərinin (AWS - Autonomous Weapon Systems) beynəlxalq hüquqi normalardan və etik çərçivələrdən kənar, qeyri-mütənasib şəkildə istifadə edilməsi riski bütün qlobal ictimaiyyət üçün həm nəzəri, həm də praktiki müstəvidə çox ciddi narahatlıqlar doğurur. Maşınların kimin yaşayıb, kimin öləcəyinə qərar verməsi beynəlxalq humanitar hüququn ən təməl prinsiplərini sınağa çəkir. Digər tərəfdən, qlobal kiberməkanda həyata keçirilən kiber hücumların artıq təxminən 40%-nin Sİ dəstəklili olması, zərərli proqramların (malware) ənənəvi antivirus sistemlərindən yayınmaq üçün öz kodlarını dinamik şəkildə dəyişdirə bilməsi (polymorphic malware) dövlətlərin kritik infrastrukturunu böyük təhlükə altına atır.

Təqdim olunan bu konfrans materialı problemin qoyuluşunu daha da dərinləşdirərək süni intellektin hərbi və kiber təhlükəsizlik sahəsindəki tətbiqlərinin riyazi, texniki, etik və strateji ölçülərini vahid bir konseptual çərçivədə analiz edir. Araşdırmada maşın öyrənməsi (ML) və dərin öyrənmə (DL) modellərinin hücum və müdafiə ssenarilərində tətbiq edilən mürəkkəb riyazi tənlilikləri, mərkəzləşdirilmiş PUA (Pilotsuz Uçuş Aparatı) sürülərinin idarə edilməsi mexanizmləri və əldə olunmuş elmi nəticələrin həm qlobal, həm də Azərbaycan Respublikasının lokal kibertəhlükəsizlik ekosistemi üçün elmi-praktiki əhəmiyyəti ətraflı şəkildə ortaya qoyulur. Tədqiqatın əsas məqsədi süni intellektin sadəcə texnoloji bir alət deyil, hərbi-siyasi reallıqları yenidən təyin edən qlobal bir fenomen olduğunu sübut etmək və bu texnologiyanın məsuliyyətli, təhlükəsiz istifadəsi üçün gərəkli olan arxitekturaları müəyyənəşdirməkdir.

II. HƏRBİ QƏRAR QƏBULETMƏ PROSESİNDƏ SÜNİ İNTELLEKT VƏ OODA DÖVRÜNÜN DİNAMİKASI

Hərbi Qərar Qəbuletmə Prosesi (HQQP) ənənəvi olaraq böyük qərargahların, çoxsaylı zabitlərin və uzun saatlar tələb

edən analitik işlərin cəmləşdiyi mürəkkəb bir planlaşdırma dövrüdür. Lakin müasir müharibənin asimmetrik və hiper-dinamik təbiəti bu ənənəvi "Napolyon dövrü" idarəetmə modelini səmərəsizləşdirmişdir. Bu nöqtədə süni intellekt HQQP-nin bütün mərhələlərini kökündən avtomatlaşdıraraq hərbi komandanlıqlara böyük strateji üstünlük qazandırır. Sİ təməlli Qərar Dəstək Sistemləri (QDS) müxtəlif mənbələrdən gələn böyük həcmli məlumatları anında emal edərək komandirlərə Fəaliyyət Tərzləri (FT) təklif edir və komandanlığın koqnitiv yükünü minimuma endirir.

OODA dövrünün sürətlənməsi və Sİ inteqrasiyası

Hərbi əməliyyatların idarə edilməsi və taktiki qərarların qəbulu baxımından mərkəzdə dayanan ən vacib nəzəri konseptlərdən biri ABŞ Hərbi Hava Qüvvələrinin polkovniki Con Boyd (John Boyd) tərəfindən irəli sürülmüş OODA dövrüdür (Observe, Orient, Decide, Act - Müşahidə, Səmti müəyyənləşdirmə, Qərar vermə və Fəaliyyət) [1]. Hərbi doktrinalara görə, rəqibindən daha sürətli OODA dövrünə malik olan tərəf təşəbbüsü ələ alır, düşməni reaktiv vəziyyətə salır və nəticədə qələbə qazanır. Süni intellekt bu dörd mərhələnin hər birində insan məhdudiyyətlərini aradan qaldıraraq idarəetməni misli görünməmiş dərəcədə optimallaşdırır (Cədvəl 1):

- **Müşahidə (Observe):** Bu mərhələ döyüş meydanından məlumatın toplanmasını ehtiva edir. Sİ sistemləri peyk görüntüləri, dronlardan gələn optik və infraqırmızı sensor məlumatları, radar siqnailləri, düşmənin şəbəkələrindən əldə edilən kibernetik telemetriya və açıq mənbəli kəşfiyyat (OSINT) hesabatları daxil olmaqla, petabaytlarla məlumatı real-vaxt rejimində mənimsəyir. Verilənlər şəbəkələri (data mesh) və paylanmış axtarış arxitekturaları vasitəsilə ən qeyri-müəyyən mühitlərdə belə dəqiq informasiya toplanır. Döyüş dumanı içərisində insan gözünün görə bilməyəcəyi mikro-hərəkətlər kompüter görməsi (computer vision) alqoritmləri tərəfindən anında təsbit edilir [2].
- **Səmti müəyyənləşdirmə (Orient):** Bu, OODA dövrünün ən mürəkkəb mərhələsidir. Burada məsələ yalnız məlumat toplamaq deyil, o məlumatın təhlili, kontekstləşdirilməsi və prioritetləşdirilməsidir. Süni intellekt yığılmış xam məlumatlar arasındakı gizli nümunələri (patterns) və anomaliyaları aşkar edir. Böyük dil modelləri (LLM) və multimodal süni intellekt sistemləri vasitəsilə bu məlumatlar komandirin anlaya biləcəyi vizual və ya mətn əsaslı qeydlərə çevrilir. Rəqibin ehtimal edilən niyyəti və zəif nöqtələri saniyələr içində modelləşdirilir [3].
- **Qərar vermə (Decide):** Ənənəvi qərarqəbularda bu proses günlər çəkə bilər. Lakin avtomatlaşdırılmış idarəetmə mərkəzləri müxtəlif ehtimalları hesablayaraq saniyələr ərzində bir neçə alternativ Fəaliyyət Tərzi (FT) formalaşdırır. Süni intellekt resursların logistik vəziyyətini, hava şəraitini, relyefi və düşmənin qüvvələrinin sıxlığını nəzərə alaraq ən optimal yolu seçir. Eksperimental olaraq istifadə edilən "COA-GPT" (Generative Pre-trained Transformers) kimi sistemlər minlərlə permutasiyanı sınaqdan keçirərək qərar qəbuledicinə ən effektiv strategiyayı təklif edir [4].

- **Fəaliyyət (Act):** Sİ təkcə qərarı tövsiyə etmir, bəzi hallarda onu birbaşa icra edərək avtonom döyüş vahidlərinə signal ötürür, artilleriya sistemlərinin koordinatlarını yeniləyir və ya kiber hücumu dərhal neytrallaşdırır.

CƏDVƏL 1. OODA DÖVRÜNDƏ İNSAN VƏ Sİ ƏSASLI YANAŞMALARIN MÜQAYİSƏSİ

OODA mərhələləri	Ənənəvi hərbi yanaşma	Sİ dəstəklə yanaşma
Müşahidə	İnsan tərəfindən idarə olunan kəşfiyyat patrulları, məhdud sensor məlumatı	Çoxdomenli (quru, hava, dəniz, kosmos, kiber) paylanmış sensor şəbəkələrinin avtomatlaşdırılmış sintezi
Səmti müəyyənləşdirmə	Zabitlərin xəritələr üzərində manual hesablama aparması, gecikmiş kəşfiyyat analizi	Dərin öyrənmə alqoritmləri ilə anomaliyaların tərfi, real-vaxtda təhdid ehtimallarının vektorlaşdırılması
Qərar vermə	İnsan intuisiyası və təcrübəsi, məhdud sayda fəaliyyət tərzinin (FT) qiymətləndirilməsi	Minlərlə FT-nin simulyasiyası, riyazi optimallaşdırma və riyazi oyunlar nəzəriyyəsi ilə ən yaxşı gedişin seçilməsi.
Fəaliyyət	Sifarişlərin silsilə vasitəsilə şəfahi və ya yazılı ötürülməsi, yüksək koordinasiya xətləri	Rəqəmsal komanda vasitəsilə maşından-maşına (M2M) birbaşa əmrlərin verilməsi, mikrosaniyəlik kiber reaksiya

Riyazi məhdudiyyətlər: Klauzevits sürtünməsi və qeyri-xəttilik

Süni intellektin HQQP-yə inteqrasiyası inqilabi üstünlüklər təmin etsə də, ciddi riyazi və konseptual məhdudiyyətlərlə üzləşir. Tədqiqatlar göstərir ki, Sİ modelləri deterministik və qaydaları əvvəlcədən bəlli olan qapalı sistemlərdə (məsələn, şahmat, Go və ya müəyyən video oyunlarında) mükəmməl fəaliyyət göstərir. Lakin müharibə cəbri və ya xətti deyil, olduqca qeyri-xətti (non-linear), xaotik və dinamik bir mühitdir. Məşhur hərbi nəzəriyyəçi Karl fon Klauzevitsin (Carl von Clausewitz) vurğuladığı "müharibə sürtünməsi" (friction in war) və "müharibənin dumanı" (fog of war) kompüter alqoritmləri tərəfindən heç vaxt tam olaraq ölçülə və proqnozlaşdırıla bilməz [5].

Eksperimental COA-GPT sisteminin qiymətləndirilməsi zamanı bu problem özünü qabarıq şəkildə büruzə vermişdir. Sistem qərarların keyfiyyətini üç əsas kəmiyyət metriki əsasında qiymətləndirir: Ümumi Mükafat (Total Reward - məsələn, ələ keçirilən ərazi üçün +10 xal, itirilən ərazi üçün -10 xal), Dost Qüvvələrin İtkiləri (Friendly Force Casualties) və Düşmənin İtkiləri (Threat Force Casualties). Lakin yalnız bu kvantitativ göstəricilərə söykənmək "həddindən artıq uyğunlaşma" (overfitting) və tarixi reallıqdan uzaqlaşma riski yaradır. Tarix göstərir ki, müharibənin sırf riyazi hesablamalarla aparılması böyük fəlakətlərə gətirib çıxarır. Məsələn, ABŞ-in Vyetnam müharibəsindəki Müdafiə Naziri Robert MakNamaranın (Robert McNamara) sırf düşmənin cəsədlərinin sayılmasına əsaslanan "ölü sayı" (body count) strategiyası cəmiyyətin iradəsini, mənafeğini və psixoloji dözümlülüyünü nəzərə almadığı üçün tamamilə uğursuz olmuşdu. Süni intellekt "Vətənpərvərlik", "Qorxu", "İnam" və "Xalqın dəstəyi" kimi qeyri-maddi, riyazi olaraq kodlaşdırılması qeyri-mümkün olan abstrakt insan amillərini hesablama bilmir [5].

Riyazi olaraq, hərbi Sİ alqoritmləri qeyri-müəyyənliyi idarə etmək üçün Bayes şəbəkələrinə (Bayesian networks), Monte Karlo simulyasiyalarına və Qismən Müşahidə Edilə Bilən Markov Qərar Proseslərinə (POMDP - Partially Observable Markov Decision Processes) müraciət etməyə çalışır. Dinamik döyüş meydanında ehtimalların paylanması POMDP modelində aşağıdakı kimi formalaşdırılır:

$$P(S_{t+1} | S_t, A_t) = \sum_{o \in O} P(S_{t+1} | S_t, A_t, o) P(o | S_t)$$

Burada S_t müəyyən bir t anında döyüş meydanının ümumi vəziyyətini, A_t tətbiq edilən Fəaliyyət Tərzini (FT), o qeyri-müəyyən xarici müşahidə faktorlarını (məsələn, düşmənin gizli manevrləri, hava şəraitinin qəfil pisləşməsi), S_{t+1} isə atılan addımdan sonrakı yeni hərbi vəziyyəti təmsil edir [6]. Baxmayaraq ki, bu riyazi modellər risk intervallarını və inam ehtimallarını hesablayır, problem ondadır ki, müharibənin mədəni, psixoloji və siyasi parametrləri Sİ sisteminin ehtimallar nəzəriyyəsinə təmiz kəmiyyət kimi daxil edilə bilmir. Buna görə də, müasir hərbi doktrinalar riyazi modellərin sırf "qüvvə planlaşdırılması" (force-planning) aləti kimi istifadə edilməsini, son və həlledici qərarın isə mütləq şəkildə insan tərəfindən (Human-in-the-Loop) verilməli olduğunu vurğulayır. İnsan mühakiməsi maşının edə bilməyəcəyi əxlaqi və strateji tarazlığı təmin edən yeganə mexanizmdir [5].

III. RƏQİB MAŞIN ÖYRƏNMƏSİ VƏ SİSTEM ZƏİFLİKLƏRİ

Sİ alqoritmlərinin hərbi mühitdəki ən böyük zəifliyi klassik proqram təminatından fərqli olaraq, intellektual rəqiblərə və bilərəkdən edilən manipulyasiyalara qarşı çox həssas olmasıdır. Bu boşluqdan istifadə edərək Sİ modellərini çaşdırmaq, aldatmaq və onları yanlış, bəzən isə fəlakətli qərarlar verməyə məcbur etmək üçün yaradılmış elm sahəsi Rəqib Maşın Öyrənməsi (Adversarial Machine Learning - AML) adlanır. AML təhdidi süni intellektin təhlükəsizliyini sarsıdan ən kəskin problemlərdən biridir, çünki o birbaşa Dərin Neyron Şəbəkələrinin (DNN) daxili struktur və öyrənmə qüsurlarından qaynaqlanır.

Hərbi təyinatlı kibermüdafiə, avtonom idarəetmə və təsvir tanıma (image recognition) sistemlərində 3 əsas AML hücum vektoru mövcuddur:

Məlumat zəhərlənməsi (Poisoning attacks): Bu hücum növü sistemin ən zəif olduğu mərhələdə – Sİ modelinin "təlim fazası" (training phase) zamanı baş verir. Düşmən kiber qüvvələri tərəfindən verilənlər bazasına xüsusi olaraq dizayn edilmiş zərərli məlumatlar, şifrələr və ya manipulyasiya edilmiş şəkillər sızdırılır. Nəticədə neyron şəbəkəsinin nümunələrin tanınması (pattern recognition) imkanı kökündən pozulur. Gələcəkdə, döyüş meydanında bu "zəhərlənmiş" model tətbiq edilərkən, Sİ dost qüvvələrin zirehli maşınlarını düşmən maşını olaraq tanıya bilər ki, bu da "dost atəşi" (friendly fire) probleminə səbəb olar [7].

Yayınma hücumları (Evasion attacks): "Sınaq və ya İstismar" (testing/deployment) fazasında tətbiq edilir. Bu zaman Sİ modelinin daxili strukturuna və ya təlim məlumatlarına heç bir dəyişiklik edilmir. Əvəzində, düşmən daxil edilən siqnalara və ya optik obyektlərə insan gözü ilə fərqi nə verilə bilməsinə imkan

olmayan, hesablanmış kiçik bir riyazi "küy" (perturbation) əlavə edir [8]. Məsələn, səmada uçan düşmən qırıcı təyyarəsinin gövdəsinə çəkilən spesifik bir naxış Sİ-nin kamerası tərəfindən qəbul edildikdə, sistemin neyronları bunu qırıcı deyil, sıradan bir quş, bulud və ya kənd təsərrüfatı obyektini kimi emal edə bilər. Belə optik illüziyalar hədəf tanıma sistemlərini tamamilə iflic edir.

Ekstraksiya hücumları (Extraction attacks): Modelin gizli daxili alqoritminin, konfidensial hərbi təlim məlumatlarının nüsxəsinin çıxarılması məqsədi daşıyır. Rəqib, Sİ sisteminə davamlı olaraq minlərlə müxtəlif sorğu göndərir və çıxış nəticələrini izləyir [7]. Bu reaksiyalara əsaslanaraq əks-mühəndislik (reverse-engineering) vasitəsilə sistemin arxitekturasının kopyası yaradılır və dövlət sirri təşkil edən hərbi məlumatlar ifşa olunur.

AML hücumlarından qorunmaq hərbi sənaye üçün həyati əhəmiyyət daşıyır. Bu məqsədlə "Rəqib Təlimi" (Adversarial Training) metodu vasitəsilə modellərə bilərəkdən manipulyasiya edilmiş şəkillər göstərilir ki, Sİ hücumlara qarşı immunitet qazansın. "Qırmızı Komanda" (Red Teaming) mütəxəssisləri davamlı olaraq öz ordularının Sİ sistemlərinə etibarlı kiber hücumlar təşkil edərək zəiflikləri aşkarlayırlar. Daha innovativ müdafiə yanaşması isə hibrid modellərin qurulmasıdır: Dərin Neyron Şəbəkələrinin (DNN) həssaslığını kompensasiya etmək məqsədilə, arxa planda Klassik Maşın Öyrənməsi (məsələn, Təsadüfi Meşə - Random Forest, və ya Dəstək Vektor Maşınları - SVM) modelləri "İkinci Dərəcəli Yoxlama" (Secondary Verification) sistemi kimi fəaliyyət göstərir. Klassik ML modelləri neyron arxitekturasına əsaslanmadığı üçün yayınma hücumlarının riyazi manipulyasiyalarına qarşı təbii olaraq toxunulmazlıq (immunity) nümayiş etdirirlər və qərar qəbul etmənin doğruluğunu təmin edirlər. NATO tədqiqatlarına görə, Sİ modellərinin etibarlılığını qorumaq üçün hərbi məlumatların ciddi yoxlanılması (data vetting) və hərbi Sİ modelinin təlim dövrü ilə aktiv əməliyyat dövrünün bir-birindən fiziki və rəqəmsal olaraq kəskin şəkildə izolyasiya edilməsi məcburi olmalıdır [5].

IV. İZAH EDİLƏ BİLƏN SÜNİ İNTELLEKT VƏ QARA QUTU PROBLEMİNİN HƏLLİ

Hərbi komandanlıq döyüş meydanında hər hansı bir avtomatlaşdırılmış qərara etibar etmək və onun hüquqi-əxlaqi məsuliyyətini daşımaq üçün, qərarların maşın tərəfindən "necə" və "niyə" verildiyini tam olaraq anlamaq məcburiyyətindədir. Müasir generativ Sİ və xüsusən çoxqatlı Dərin Öyrənmə alqoritmləri hərbi dairələrdə "qara qutu" (black-box) problemi yaradır [5]. Məsələn, kompleks bir neyron şəbəkəsi qərar qəbul edərkən on milyonlarla, bəzən milyardlarla parametrlər və çəki əməli üzərindən əməliyyat aparır. Əgər Sİ düşmən mövqeyini deyil, səhvən mülki obyekti hədəf olaraq seçirsə, arxada çalışan alqoritmin hansı hesablamaya əsasən bu qərara gəldiyi nə komandır, nə də proqramçı üçün aydın olmur.

İnsan və maşın koalisiyasındakı bu inam boşluğunu doldurmaq üçün son illərdə İzah Edilə Bilən Süni İntellekt (Explainable AI - XAI) arxitekturası hərbi standartlara uyğunlaşdırılaraq tətbiq olunmağa başlanmışdır. XAI-nin əsas məqsədi mürəkkəb riyazi alqoritmlərin nəticələrini insanın anlaya biləcəyi dildə izah etmək və modelin daxili məntiqini

şəffaflaşdırmaqdır. Xüsusilə, hərbi kiber əməliyyatlarda və anomaliyaların aşkarlanmasında LIME (Local Interpretable Model-Agnostic Explanations) və SHAP (SHapley Additive exPlanations) kimi fərqli riyazi interpretasiya modellərindən istifadə edilir. Bu izahat modelləri hərbi qərarların xüsusiyyətlər səviyyəsində (feature-level) aydınlığını artıraraq, qərara təsir edən ən vacib amilləri sıralayır [9].

XAI-nin hərbi kiber əməliyyatlarda tətbiqi İnsan-Mərkəzli İzah Edilə Bilən Sİ (HCXAI) konsepti daxilində həyata keçirilir. Lakin burada çox həssas bir dilemma ortaya çıxır: Əməliyyat təhlükəsizliyi (OPSEC) ilə Şəffaflıq arasındakı ziddiyyət. Əgər Sİ algoritmi çox şəffaf və izah edilə bilən olsa, düşmən qüvvələri həmin modelin daxili işləmə məntiqini daha asanlıqla öyrənərək Rəqib Maşın Öyrənməsi (AML) vasitəsilə sistemi manipulyasiya edə bilər. Buna görə də, hərbi sahədə kriptografik cəhətdən təmin edilmiş XAI mexanizmlərinə ehtiyac duyulur. Ümumi olaraq, XAI tətbiqləri komandirlərə qüsurları ilkin mərhələdə tapmağa kömək edir, beynəlxalq hüquqi normalara uyğunluğu yoxlayır, sistemin qərəzliliyini aradan qaldırır və müharibə zonalarında mülki şəxslərə dəyə biləcək yanlış zərərlərin (collateral damage) qarşısının alınması üçün həyati əhəmiyyət daşıyır [10].

V. BEYNƏLXALQ HÜQUQ, ETİK ÇƏRÇİVƏLƏR VƏ İNSAN HÜQUQLARI

Süni intellektin sürətlə hərbiləşdirilməsi təkcə texniki və texnoloji tərəqqi məsələsi deyil, o eyni zamanda global miqyasda görünməmiş etik, mənəvi və beynəlxalq hüquqi debatlar yaratmışdır. Birləşmiş Millətlər Təşkilatının (BMT) hesabatlarına, eləcə də Beynəlxalq Qırmızı Xaç Komitəsinin (ICRC) təqdim etdiyi rəsmi sənədlərə əsasən, maşınlar insanın həyatı və ölümü üzərində hökm vermək, hədəfi seçib müstəqil şəkildə yox etmək səlahiyyətinin verilməsi Beynəlxalq Humanitar Hüququn (IHL) və Beynəlxalq İnsan Hüquqları Hüququnun (IHRL) fundamental dayaqlarını - "proporsionallıq" (proportionality), "hədəfi ayırd etmə" (distinction) və "məsuliyyət/hesabatlılıq" prinsiplərini birbaşa risk altına salır [11].

Etik risklərin nə qədər ciddi olduğu yaxın keçmişdə baş vermiş real müharibə ssenarilərində özünü sübut etmişdir. Məsələn, bəzi tədqiqatlarda bildirilir ki, münaqişə zonalarında hərbi əməliyyatların gedişində qərar vermə prosesinin avtomatlaşdırılması üçün "Lavender" kimi adlandırılan Sİ proqramlarından istifadə edilmişdir. Təqdim olunan qeyri-rəsmi araşdırmalara görə, bəzən Sİ sistemləri tərəfindən yaradılan və 37.000-dən çox potensial hədəfi siyahıya alan avtomatlaşdırılmış hesabatlar komandanlığa təqdim olunmuş, hərbi personal tərəfindən bir hədəfin təsdiqlənməsi üçün insan nəzarətinə cəmi 20 saniyə vaxt ayrılmışdır. Maşının çıxardığı nəticələrdəki xəta payının 10% ətrafında olmasına baxmayaraq, yüksək sürətə olan ehtiyac operatorların alqoritmə kor-koranə inanmasına, yəni avtomatlaşdırma qərəzliliyinə (automation bias) məruz qalmasına səbəb olmuşdur. Alqoritmlər hədəfləri çox vaxt onların ən müdafiəsiz olduğu məkanlarda (evlərində, ailələrinin yanında) nişan aldığından, insan faktoru və proporsionallıq mülahizələri arxa plana keçmiş, mülki itkilər barədə kəskin hüquqi tənqidlər ortaya çıxmışdır [12].

Həmçinin, Sİ alqoritmlərinin tarixi məlumatlarla təlimləndirilməsi zamanı yol verilən yanlışlıqlar və alqoritmik qərəzlilik, Mülki və Siyasi Hüquqlar haqqında Beynəlxalq Paktın (ICCPR) və Qadınlara qarşı ayrı-seçkiliyin bütün formalarının ləğv edilməsi haqqında Konvensiyanın (CEDAW) prinsiplərini poza bilər. Qərəzli sistemlər müəyyən demografik qrupları, xüsusən qadın və uşaqlar kimi həssas şəxsləri hərbi hədəflərlə səhv sala bilər. Hərbi etika mütəxəssislərinin fikrincə, davamlı olaraq maşınlar tabe olan insan operatorlarında əxlaqi və dəyər-mərkəzli düşüncə tərzini körelir, onlar sadəcə təsdiqləyici bir monitoring cihazına çevrilir və "mənəvi tamaşaçı" (moral spectator) statusu qazanırlar. Bu global təhlükənin qarşısının alınması üçün bütün dünyada BMT və digər təşkilatlar tərəfindən qəbul edilməsi təklif olunan normativ sənədlər təkid edir ki, Silahlı Qüvvələr bütün avtomatlaşdırılmış "sensor-dan-aticıya" (sensor-to-shooter) sistemlərində və öldürücü silahlarda "Mənalı İnsan Nəzarətini" (Meaningful Human Control) hüquqi öhdəlik kimi qoruyub saxlamalıdır.

VI. MİLLİ TƏHLÜKƏSİZLİK STRATEGİYALARI VƏ AZƏRBAYCAN KONTEKSTİ

Tarixi və geosiyasi reallıqlar sübut edir ki, pilotsuz uçuş aparatları, süni intellekt və kibermüharibə elementləri artıq gələcəyin deyil, bugünkü müasir konfliktlərin asimmetrik mahiyyətini müəyyənləşdirən əsas amillərdir. Qlobal güclərin Sİ sahəsindəki kəşfləri yeni bir Sİ silahlanma yarışına (AI arms race) çevirdiyi indiki tarixi dövrdə, Azərbaycan Respublikası da regional reallıqları nəzərə alaraq öz milli təhlükəsizlik konsepsiyasını global kiber və Sİ təhdidlərinə sürətlə uyğunlaşdırmaqdadır.

Azərbaycanın bu sahədəki milli yanaşması və təcrübəsi bir neçə mühüm strateji sütuna əsaslanır:

Döyüş meydanında avtomatlaşdırma və innovativ müharibə təcrübəsi:

Yaxın keçmişdə baş vermiş regional hərbi əməliyyatlar, xüsusən də 2020-ci il Vətən Müharibəsi global hərbi ekspertlər tərəfindən innovativ müharibə sənətinin ən parlaq nümunəsi kimi göstərilir. Müstəqil tədqiqat və analizlərin nəticələri göstərir ki, hərbi əməliyyatların gedişində yüksək texnologiyalı avtomatlaşdırılmış idarəetmə sistemlərinin və PUA-ların uğurlu tətbiqi müharibənin taleyini həlledici dərəcədə təyin etmiş, məhv edilən hədəflərin təqribən 45%-i məhz bu platformalar vasitəsilə neytrallaşdırılmışdır. Kəşfiyyat (ISR), logistika, qərar qəbul etmə və birbaşa vuruş sistemlərinin avtomatlaşdırılması ənənəvi ordu strukturlarının çevikliyini inanılmaz dərəcədə artırmış, komanda-nəzarət zəncirini optimallaşdırmışdır [13].

Milli kibertəhlükəsizlik ekosisteminin qurulması:

Qlobal informasiya məkanında rəqəmsal sərhədlərin qorunması məqsədilə ölkə rəhbərliyi tərəfindən davamlı islahatlar həyata keçirilir. Prezident İlham Əliyevin müvafiq fərmanı ilə 2023-cü il 28 avqust tarixində təsdiqlənmiş "Azərbaycan Respublikasının informasiya təhlükəsizliyi və kibertəhlükəsizliyə dair 2023–2027-ci illər üçün Strategiyası" dövlətin rəqəmsal suverenliyini təmin etmək üçün dərin bir yol xəritəsi ortaya qoyur. Bu çoxşaxəli strategiya çərçivəsində Milli Təhlükəsizlik Əməliyyatları Mərkəzi (N-SOC) və Milli

Kompüter İnsidentlərinə Qarşı Mübarizə Mərkəzinin (N-CERT) fəaliyyətə başlaması nəzərdə tutulur. Bu təsisatlar kiber insidentlərin global kiber-kəşfiyyat vasitələri və Sİ metodologiyaları ilə real-vaxt rejimində monitorinq edilməsi, analiz olunması və dərhal qarşısının alınması üçün unikal bir dövlət mərkəzi rolunu oynayır.

Hüquqi və kadr dayanıqlılığının gücləndirilməsi:

2025-ci ildə elan edilmiş "Rəqəmsal İnkişaf Konsepsiyası"na əsasən, informasiya resurslarının konsolidasiyası, ixtisaslaşmış kadr potensialının formalaşdırılması, beynəlxalq əməkdaşlıq çərçivəsində "Azərbaycan Kibertəhlükəsizlik Mərkəzi"nin fəaliyyətinin genişləndirilməsi (qabaqcıl mütəxəssislərin yetişdirilməsi) kimi əhəmiyyətli addımlar atılmışdır.

Kritik informasiya infrastrukturunun (Kİİ) Sİ təməlli təhdidlərdən qorunması məqsədilə kiber gigiyena (Cyber Hygiene) layihələri, dövlət sistemlərinin məcburi informasiya auditi qaydaları tətbiq olunur. Bunun nəticəsidir ki, beynəlxalq "Global Cybersecurity Index" qiymətləndirməsində ölkə öz hüquqi, texniki və təşkilati göstəricilərinə görə yüksək sıralarda qərarlaşmışdır [14].

Bütün bu dövlət proqramları göstərir ki, Sİ və kiber hücumların hərbiləşdirilməsi kimi yeni nəsil təhdidlərə qarşı Azərbaycan Respublikası passiv müdafiə mövqeyində deyil, proaktiv, institusional və qanunvericilik istiqamətində qabaqlayıcı addımlar atır, ən son texnologiyalarla hərbi və informasiya qüdrətini davamlı balanslaşdırır.

NƏTİCƏ

Həyata keçirilmiş bu genişmiqyaslı tədqiqatın yekunları sübut edir ki, süni intellektin hərbiləşdirilməsi yalnız texnoloji alətin inkişafı deyil, əksinə, global milli təhlükəsizlik konsepsiyalarını, beynəlxalq münasibətlər sistemini, ənənəvi hərbi doktrinaları və qanunvericilik normalarını kökündən dəyişdirən fundamental fəlsəfi və strateji paradigma dəyişikliyidir.

Tədqiqat nəticəsində aydın olur ki, maşın öyrənməsi və qabaqcıl idarəetmə sistemləri HQQP-nin "OODA dövrünü" misli görünməmiş dərəcədə sürətləndirərək komandanlıq iyerarxiyasını daha effektiv formaya salır. Bundan əlavə, tədqiqat dərin neyron şəbəkələrinin (DNN) kiberməkanda Rəqib Maşın Öyrənməsi (AML), yayınma (evasion) və məlumat zəhərlənməsi (poisoning) hücumlarına qarşı daxili struktur zəifliklərini riyazi əsaslarla aşkara çıxarmış, bu risklərin qarşısını almaq üçün Klassik Maşın Öyrənməsi və İzah Edilə Bilən Sİ (XAI) arxitekturalarının necə hibridləşdirilməsi zərurətini elmi sübutlarla əsaslandırmışdır.

Əldə olunmuş nəticələr silahlı qüvvələrin strukturları, hökumət siyasətinin tərtibatçıları və milli kiber təhlükəsizlik agentlikləri üçün birbaşa tətbiqi və strateji əhəmiyyət kəsb edir. Birincisi, HQQP-də Sİ-nin riyazi məhdudiyyətləri (insan emosiyalarını, xaosu və əxlaqi hesablaya bilməməsi) dərk edilməli və avtomatlaşdırma qərəzliliyindən qaçınmaq məqsədilə həlledici düymənin başında "Mənalı İnsan Nəzarəti" (Meaningful Human Control) hüquqi imperativ kimi qorunmalıdır.

İkincisi, dövlətlərin artan "Sİ hücum səthinə" və kiber cinayətkarlığa qarşı öz daxili rəqəmsal və qanunvericilik immunitetini necə artırmalı olduğuna dair Azərbaycan Respublikasının 2023-2027-ci illər Kibertəhlükəsizlik Strategiyası və formalaşdırılan Kiber Mərkəzlər (N-SOC, N-CERT) praktiki istinad nöqtəsi rolunu oynayır. Gələcək milli doktrinaların və kibermüdafiə ssenarilərinin yaradılmasında Süni İntellektə sadəcə düşməni yox edəcək bir silah kimi deyil, eyni zamanda global sabitliyi qoruyan, davamlı şəbəkə anomaliyalarını izləyən, etik və hüquqi standartlara cavab verən bir "intellektual müdafiə qalxanı" kimi yanaşılmalıdır.

ƏDƏBİYYAT

- [1] L. MacVittie, "AI and the OODA loop: Reimagining operations," F5, May 27, 2024. <https://www.f5.com/company/blog/ai-and-the-ooda-loop-reimagining-operations>
- [2] C. McEwen and O. Abramovych, "AI-ready defence: Accelerating OODA with unified search and AI," Elastic, Jan. 21, 2026. <https://www.elastic.co/blog/seach-and-ai-accelerating-defence-ooda-loop>
- [3] T. Yasuno, "RAPTOR-AI for disaster OODA loop: Hierarchical multimodal RAG with experience-driven agentic decision-making," arXiv preprint arXiv:2602.00030, 2026. <https://doi.org/10.48550/arXiv.2602.00030>
- [4] J. N. Adler, "Modernizing military decision-making: Integrating AI into Army planning," Military Review, Aug. 2025. <https://www.armyupress.army.mil/Journals/Military-Review/Online-Exclusive/2025-OLE/Modernizing-Military-Decision-Making/>
- [5] V. Gallitelli, "Artificial intelligence-enabled military decision-making process: The forgotten lessons on the nature of war," Journal of Advanced Military Studies, vol. 16, no. 2, 2025, doi: 10.21140/mcu.20251602005.
- [6] P. Poupard, "Partially observable Markov decision processes," in Encyclopaedia of Machine Learning, C. Sammut and G. I. Webb, Eds. Springer, 2011, doi: 10.1007/978-0-387-30164-8_629.
- [7] S. Arif, "Understanding adversarial AI: The military lens," The Forge, Dec. 8, 2025. <https://theforge.defence.gov.au/article/understanding-adversarial-ai-military-lens>
- [8] E. Alhajjar, "Adversarial machine learning poses a new threat to national security," SIGNAL Magazine, Jul. 1, 2022. <https://www.afcea.org/signal-media/cyber-edge/adversarial-machine-learning-poses-new-threat-national-security>
- [9] I. Syllaidopoulos, K. S. Ntalianis, and I. Salmon, "A comprehensive survey on AI in counter-terrorism and cybersecurity: Challenges and ethical dimensions," IEEE Access, vol. 13, pp. 91740–91764, 2025, doi: 10.1109/ACCESS.2025.3572348.
- [10] D. Staněk, I. Klaban, and A. Coufalíková, "Explainable artificial intelligence: State of the art and beyond," in Proc. 2025 Int. Conf. Military Technol. (ICMT), 2025, pp. 1–7, doi: 10.1109/ICMT11061287.
- [11] C. Droege, "UN Security Council: We cannot let AI be deployed on the battlefield without oversight and regulation," International Committee of the Red Cross, Sep. 26, 2025. <https://www.icrc.org/en/statement/we-cannot-let-AI-be-deployed-on-battlefield-without-oversight-and-regulation>
- [12] A. A. Satti, "AI: The new paradigm of war," Annenberg School for Communication, 2026. <https://www.asc.upenn.edu/research/centers/milton-wolf-seminar-media-and-diplomacy-8>
- [13] S. M. Babayev, "The role of artificial intelligence in automating military operations and enhancing combat readiness," in Coll. Sci. Papers Works Materials IV Int. Sci. Conf.: Global Challenges and Innovations: Ways of Developing Modern Science, UKRLOGOS Group, 2025, pp. 107–108,
- [14] International Trade Administration, "Azerbaijan - Information and communications technology," U.S. Department of Commerce, Jan. 12, 2026. <https://www.trade.gov/country-commercial-guides/azerbaijan-information-and-communications-technology>

Weaponization of Artificial Intelligence: Global Security, Mathematical Modelling and Cyber Defence Strategies

Nuran Mahmudov¹, Mehrac Huseynov²

^{1,2}National Defense University, Baku, Azerbaijan

¹mahmudovnuran@outlook.com, ²mehrac77@gmail.com

Abstract- This research material is a comprehensive report presenting a multidimensional analysis of the weaponization of artificial intelligence and its transformative impact on global security systems. The core of this study focuses on optimizing the Military Decision-Making Process and the OODA loop.

Cyber threats executed via adversarial machine learning, explainable AI architectures, automation bias, and the meaningful human control concept within international humanitarian law are thoroughly evaluated. Finally, the national cybersecurity strategy of Azerbaijan, operations centers and digital sovereignty are explored.

Keywords- artificial intelligence; weaponization; cybersecurity; autonomous weapons systems; machine learning.