

Նարեկ Մ. Պողոսյան
պատմական գիտությունների թեկնածու

**ՑԵՂԱՍՊԱՆՈՒԹՅԱՆ ԵՎ ԶԱՆԳՎԱԾԱՅԻՆ ՍՊԱՆՈՒԹՅՈՒՆՆԵՐԻ
ԻՐԱԳՈՐԾՄԱՆ ՀՆԱՐԱՎՈՐ ՎՏԱՆԳԸ ԱՐՀԵՍՏԱԿԱՆ ԲԱՆԱԿԱՆՈՒԹՅԱՆ
ԶԱՐԳԱՑՄԱՆ ՀԱՄԱՏԵՔՍՏՈՒՄ¹**

Բանալի բառեր՝ արհեստական բանականություն, մահացու ինքնավար գեներ, ոռոտներ, ցեղասպանություն, խմբերի թիրախավորում, սոցիալական մեդիա, կեղծ տեսանկյուն, արհեստական բանականության էթիկա:

Տեխնոլոգիաների արագ զարգացումը, մասնավորապես՝ արհեստական բանականության կիրառումը, վերջին տարիներին փորձագետների մոտ լուրջ անհանգստություն է առաջ բերել. մեքենաների ինքնուրույն մտածելու և որոշումներ կայացնելու կարողությունն իր դրական կողմերով հանդերձ՝ մարդկության համար բազմաթիվ վտանգներ է ստեղծում: Արհեստական բանականության զարգացումը և ինքնավար գեները կարող են օգտագործվել ազգային, կրոնական, ռասայական որոշակի խմբերի թիրախավորման նպատակով, ինչի արդյունքում էլ մեծանում է ցեղասպանության և զանգվածային սպանությունների իրագործման վտանգը:

Ներկայում, երբ արդեն արհեստական բանականությամբ օժտված սպառազինությունների գլոբալ մրցավազքն իրողություն է, տեխնոլոգիական ոլորտի բազմաթիվ ներկայացուցիչների ուշադրությանն է արժանանում առաջին հերթին մահացու ինքնավար գեների արգելման հարցը: Բացի այդ՝ արհեստական բանականության զարգացման բացասական հնարավոր հետևանքները կանխելու համար ոլորտի մասնագետների կողմից առաջ է քաշվել արհեստական բանականության էթիկայի մշակման անհրաժեշտությունը:

Մարդկության համար արդեն իսկ հրատապ է դարձել ոռոտատեխնիկայի և արհեստական բանականության տեխնոլոգիայի զարգացման արդյունքում առաջ եկած հնարավոր վտանգը կանխելու, մարդու իրավունքների պաշտպանության և էթիկայի նորմերին համապատասխան օրենսդրական մեխանիզմների ստեղծումը:

Հաշվի առնելով այն հանգամանքը, որ արհեստական բանականությունն օգտագործվում է նաև ամենատարբեր հանցագործությունների հայտնաբերման և կանխարգելման համար, հետևաբար նպատակահարմար է այն ծառայեցնել նաև ցեղասպանության և զանգվածային սպանությունների կանխարգելմանը:

Ներկայում տեխնոլոգիական առաջընթացը և նորարարությունը մեծ փոփոխություններ են առաջացրել կյանքի գրեթե բոլոր բնագավառներում. ժամանակակից աշխարհն առանց տեխնիկական միջոցների կիրառման անհնար է պատկերացնել: Արհեստական բանականության կիրառությունը հեղաշրջիչ նշանակություն կարող է ունենալ մարդկանց և պետությունների համար՝ մի կողմից հեշտացնելով և ավել-

1 Հոդվածը ստացվել է 18.09.2019 և ընդունվել տպագրության՝ 12.01.2020:

լի հարմարավետ դարձնելով կյանքը, մյուս կողմից՝ հանգեցնելով բազմաթիվ վտանգների: Ավելին՝ արհեստական բանականությունն ու ռոբոտատեխնիկան կարող են ինչպես միլիոնավոր աշխատատեղերի կրճատման, այնպես էլ՝ մահացու ինքնավար զենքերի կիրառման պատճառ դառնալ:

Մեր օրերում աշխարհի տարբեր երկրներում նկատվում է էթնիկ, կրոնական, ռասայական գործոններով պայմանավորված թշնամանքի աճ. մեծ է վտանգը, որ արհեստական բանականության զարգացումը և ինքնավար զենքերը կարող են օգտագործվել ազգային, կրոնական, ռասայական որոշակի խմբերի թիրախավորման նպատակով, քանի որ արհեստական բանականությամբ օժտված ինքնավար զենքերը հնարավորություն կտան շատ արագ նույնականացնել այդ խմբերին պատկանող մարդկանց և անհրաժեշտության դեպքում ոչնչացնել:

Փաստորեն, ցեղասպանության իրագործման վտանգի ուժգնացմանը զուգընթաց կատարելագործվում է նաև տեխնիկան: Եթե 20-րդ դարի սկզբում Հայոց ցեղասպանության տարիներին ռճրագործներն անձամբ էին սպանում զոհերին, իսկ Հոլոքոստի ժամանակ նացիստները օգտագործում էին գազակացիկները, ապա 21-րդ դարում արհեստական բանականության կիրառությունը կարող է այնպիսի պայմաններ ստեղծել, որ մարդն ընդհանրապես մասնակցություն չունենա ցեղասպանության իրագործման բուն գործընթացին: Առանձին անհատներ, իշխանության ներկայացուցիչներ կամ ահաբեկչական խմբեր, տիրապետելով այդպիսի տեխնոլոգիաների, կարող են հեռակառավարման կամ համապատասխան ծրագրի ներմուծման միջոցով և առանց մարդկային ռեսուրսի կիրառման թիրախավորել տարբեր խմբերի և կարճ ժամանակամիջոցում ոչնչացնել նրանց: Բացառված չէ նաև այն տարբերակը, երբ ինքնավար զենքերը և ռոբոտները ինչ-որ սխալի կամ ինքնագործունեության պարագայում թիրախավորեն մարդկանց: Ուստի ինքնավար զենքերի կամ ռոբոտների կողմից ցեղասպանության կամ զանգվածային սպանությունների իրագործման դեպքում հարկ կլինի հասկանալ, թե ովքեր են պատասխանատու դրանց գործունեության համար: Հետևաբար, ցեղասպանության և մարդկության դեմ ուղղված հանցագործությունների վաղ կանխարգելման համատեքստում հրամայական է դառնում արհեստական բանականության կիրառման մեխանիզմների ստեղծումը, որի միջոցով հնարավոր կլինի սահմանափակել մարդկանց թիրախավորելու նպատակով արհեստական բանականությամբ օժտված տեխնոլոգիաների կիրառումը և սահմանել հստակ չափորոշիչներ, որը թույլ կտա կանխել դրանց կիրառման հնարավոր վտանգները:

Արհեստական բանականության զարգացումը և դրա վտանգները մարդկության համար

Արհեստական բանականության գաղափարը հայտնի է դեռևս անտիկ ժամանակաշրջանից, երբ փիլիսոփաները խորհում էին արհեստական էակների հնարավորության մասին: Նրանք մտածում էին, որ մարդկային բանականությունը մի համա-

կարգ է, որը կարող է ինքնավերացվել և կիրառվել արհեստական էակների կողմից: Արհեստական բանականության գաղափարը տեղ է գտել նաև հին հունական, չինական, հռոմեական, հնդկական և հրեական առասպելներում²:

Արհեստական բանականության հիմնադիրներից գաղտնավերլուծող և ինֆորմատիկայի մասնագետ Ալան Թյուրինգն իր «Հաշվողական տեխնիկա և բանականություն» (1950 թ.) հոդվածում ներկայացնում է մեքենայի՝ մտածելու կարողության վերաբերյալ մշակված մի թեստ: Թյուրինգի նպատակն էր թեստի միջոցով համեմատական վերլուծության ենթարկել մարդու վարքն ու արհեստական բանականությունը: Հեղինակը հանգում է այն եզրակացության, որ եթե մեքենաները կարողանան «հիմարեցնել» մարդկանց՝ ներկայանալով, թե իբր իրենք մարդ են, այդպիսով կարող են ցույց տալ իրենց մտածելու ունակությունը³:

Արհեստական բանականության առաջին համակարգիչը «Մանչեսթեր Ֆերանտի»-ն է⁴, որն ինքնուրույն և անշտապ շաշկի ու շախմատ է խաղում՝ ուսումնասիրելով հազարավոր քայլեր և գտնելով ամենատարբեր լուծումներ⁵:

1956 թ. Դարտմուտի կոնֆերանսի ժամանակ ինֆորմատիկայի մասնագետ, ամերիկացի Ջոն Մրքքարթին հնչեցնում է «արհեստական բանականություն» արտահայտությունը: Կոնֆերանսից հետո 1957-1974 թթ. մեծ առաջընթաց է գրանցվում արհեստական բանականության ոլորտում. համակարգիչները սովորում են անգլերեն խոսել, հանրահաշվական խնդիրներ լուծել, ապացուցել երկրաչափության թեորեմները և այլն⁶:

1967 թ. Ճապոնիայում Վասեդայի համալսարանն սկսում է «WABOT» նախագիծը և 1972 թ. ավարտում «WABOT-1» անունը կրող ռոբոտի ստեղծման աշխատանքները: Այն աշխարհում առաջին բանական ռոբոտն էր, որը քայլում էր, առարկաներ տեղափոխում, մարդկանց հետ ճապոներեն խոսում⁷:

1970-ական թթ. արհեստական բանականության ոլորտում հետազոտությունները քննադատության են արժանանում: Պատճառն այն էր, որ արհեստական բանականության ոլորտի հետազոտողների ակնկալիքները չափազանց մեծ էին, սակայն նրանք չկարողացան խոստացված արդյունքները իրականություն դարձնել՝ հետազոտությունների համար տրամադրվող ոչ բավարար ֆինանսական միջոցների

2 Adrienne Mayor, *Gods and Robots: Myths, Machines, and Ancient Dreams of Technology* (Princeton: Princeton University Press, 2018), 304; Donald J. Norris, *Beginning Artificial Intelligence with the Raspberry Pi* (Barrington: Apress, 2017), 1-2.

3 Alan Turing, "Computing Machinery and Intelligence," *Mind* LIX, no. 236 (1950): 433-460.

4 «Ferranti Mark 1» մեքենան ստեղծել է Ferranti համակարգչային ընկերությունը 1951 թ.:

5 Tim Jones, *Artificial Intelligence: A Systems Approach: A Systems Approach* (Boston, Toronto, London, Singapore: Jones & Bartlett Learning, 2015), 4-5.

6 Mark Skilton, Felix Hovsepian, *The 4th Industrial Revolution: Responding to the Impact of Artificial Intelligence on Business* (Cham: Springer, 2017), 76.

7 Christopher McFadden, "15 Engineers and Their Inventions That Defined Robotics," *Interesting Engineering*, July 06, 2018, <https://interestingengineering.com/15-engineers-and-their-inventions-that-defined-robotics>, դիտվել է 07.08.2019:

պատճառով: Այս ժամանակաշրջանը հայտնի է «արհեստական բանականության ձմեռ անունով» (1974-1980 թթ.)⁸:

1980-ական թթ. այս ոլորտում հետազոտություններ իրականացնելու ուղղությամբ ֆինանսական միջոցների տրամադրման ակտիվություն է նկատվում: Այս գործընթացին 1981 թ. զգալի աջակցություն է ցուցաբերում Ճապոնիայի կառավարությունը: Կարճ ժամանակ անց այլ երկրներ ևս ֆինանսական աջակցությամբ միանում են գործընթացին. սպասումները, սակայն, չեն արդարանում և սկսվում է «արհեստական բանականության երկրորդ ձմեռը» (1987-1993 թթ.): 1993-2011 թթ. հաջողություն է գրանցվում՝ պայմանավորված ֆինանսական ներդրումներով:

2012 թ. մինչև այսօրվա տվյալների մատչելիությունը, կապի և հաշվարկային հզորությունների ավելացումը բեկումնային նշանակություն ունեցան մեքենայական ուսուցման (Machine learning), նեյրոնային ցանցերի և խորքային ուսուցման (Deep learning) բնագավառներում⁹: Ըստ կանխատեսումների՝ 2030 թ. արհեստական բանականությունը 15.7 տրիլիոն դոլարի եկամուտ կապահովի համաշխարհային տնտեսությանը¹⁰:

Ելնելով իր գործառնությունից՝ արհեստական բանականությունը բաժանվում է երեք խմբի՝ թույլ կամ նեղ, ընդհանուր կամ ուժեղ և արհեստական գերբանականություն:

Արհեստական թույլ բանականությունը ծրագրավորված է որոշակի առաջադրանքներ իրականացնելու համար: Այն կիրառվում է հատուկ խնդրի լուծման համար և չի գործում դրա սահմաններից դուրս: Ներկայում տիրապետող է արհեստական թույլ բանականությունը (անձնական օգնականները՝ Apple ընկերության Siri-ն, IBM ընկերության Watson-ը, Google ընկերության Google Assistant-ը, Amazon ընկերության Alexa-ն և Microsoft ընկերության Cortana-ն, երթևեկության լույսերի կառավարման համակարգերը մի շարք չինական քաղաքներում, չաթոտերը և այլն):

Արհեստական ուժեղ բանականությունը կարող է կատարել առաջադրանքներ, ինչպես բանական էակը: Ի տարբերություն արհեստական թույլ բանականության՝ այն չի սահմանափակվում հատուկ առաջադրանքով և կարող է կատարել ցանկացած առաջադրանք, ինչպես և մարդը (օրինակ՝ Սոֆյա ռոբոտը):

Արհեստական գերբանականությունն արհեստական բանականության հիպոթետիկ ենթատեսակ է, որն ամենախելացի մարդուց անգամ խելացի է:

Բրիտանացի մաթեմատիկոս Իրվինգ Ջոն Գուդը նշում է, որ գերբանական ցանկացած մեքենա (ultraintelligent machine) կարող է գերազանցել մարդկային մտավոր գործողությունները: Գերբանական մեքենաները կարող են նախագծել ավելի հզոր մեքենաներ, որի դեպքում տեղի է ունենալու «բանականության պայթյուն», և մարդկային բանականությունը շատ հեռ է մնալու: Այդպիսով՝ գերբանական

8 Tony Boobier, *Advanced Analytics and AI: Impact, Implementation, and the Future of Work* (Chichester: John Wiley & Sons, 2018), 41.

9 *WIPO Technology Trends 2019 – Artificial Intelligence* (Geneva: WIPO, 2019), 19.

10 Anand S. Rao, Gerard Verweij, *Sizing the prize What's the real value of AI for your business and how can you capitalise?* <https://www.pwc.com/gx/en/issues/analytics/assets/pwc-ai-analysis-sizing-the-prize-report.pdf>, դիտվել է 17.09.2019:

մեքենան վերջին գյուտն է, որը մարդը երբևէ պետք է կատարի¹¹: Իսկ ամերիկացի մաթեմատիկոս և գիտաֆանտաստ գրող Վեռնոր Վինջը ժողովրդականացրեց «տեխնոլոգիական սինգուլյարություն» (technological singularity) հասկացությունը իր՝ 1993 թ. «Գալիլեո տեխնոլոգիական սինգուլյարություն» ակնարկում, որում գրում է, թե տեխնոլոգիական սինգուլյարությունը նշանակելու է մարդկային դարաշրջանի ավարտը, քանի որ նոր սուպերբանականությունը կշարունակի արդիականացնել ինքն իրեն և առաջ կընթանա տեխնոլոգիապես անհասկանալի տեմպերով¹²: Տեխնոլոգիական սինգուլյարությունը ապագայի հիպոթետիկ ժամանակաշրջան է, երբ տեխնոլոգիական աճը կդառնա անվերահսկելի և անշրջելի, ինչը կհանգեցնի մարդկային քաղաքակրթության անհասկանալի փոփոխությունների: Ամերիկացի նշանավոր գյուտարար և ապագայաբան Ռեյմոնդ Կուրցվեյլը կանխատեսում է, որ մեր մոլորակը տեխնոլոգիական սինգուլյարության է հասնելու 2035 - 2045 թթ. ընկած ժամանակաշրջանում¹³: Ըստ Կուրցվեյլի՝ սինգուլյարությունը մարդկանց թույլ կտա հաղթահարել կենսաբանական մարմինների և ուղեղի սահմանափակումները, կընդլայնվի մարդկային մտածողությունը, սակայն այն իր հետ կբերի նաև վտանգներ՝ մասնավորապես կրաձրացնի մարդկանց՝ կործանարար հակումներին համապատասխան գործելու ունակությունը¹⁴:

Տեխնոլոգիական սինգուլյարության ի հայտ գալը ենթադրում է արհեստական բանականությամբ օժտված գերհզոր ուղեղի ստեղծում, և մասնագետների մեջ մտավախություն կա, որ արհեստական բանականությունը այնքան կգարգանա ու կհեռանա մարդկային ուղեղի գործառնությունից, որ մարդկության նկատմամբ ցեղասպանություն կիրագործի¹⁵:

Ֆիզիկոս Սթիվեն Հոփհոլդը, Microsoft ընկերության հիմնադիր Բիլ Գեյթսն ու Tesla և SpaceX ընկերությունների հիմնադիր Իլոն Մասկը բազմիցս իրենց մտահոգություն են հայտնել արհեստական բանականության առաջ բերած ռիսկերի վերաբերյալ: Նրանց մտահոգությունները նույնիսկ հասնում են այնտեղ, թե արհեստական բանականության զարգացումը կարող է վտանգ ստեղծել անգամ մարդկության գոյության համար: Բիլ Գեյթսը գտնում է, որ արհեստական բանականության սպառնալիքը պետք է անհանգստացնի մարդկանց¹⁶: Նա արհեստական բանականությունը համեմատում

11 Irving John Good, *Speculations Concerning the First Ultra intelligent Machine* (Trinity College, Oxford and Atlas Computer Laboratory, Chilton, Berkshire: Elsevier, 1965), 33.

12 Vinge Vernor, "The Coming Technological Singularity: How to Survive in the Post-Human Era," in *Vision-21: Interdisciplinary Science and Engineering in the Era of Cyberspace*, ed. Geoffrey Landis (NASA Publication CP-10129, 1993), 11-22, <https://mindstalk.net/vinge/vinge-sing.html>, դիտվել է 04.05.2019:

13 Charles Jennings, *Artificial Intelligence: Rise of the Lightspeed Learners* (Lanham, London: Rowman & Littlefield, 2019), 32.

14 Ray Kurzweil, *The Singularity Is Near: When Humans Transcend Biology* (London: Penguin Books, 2005), 25, 34.

15 "Is Technological Evolution Infinite or Finite?" October 15, 2013, <https://futurism.com/is-technological-evolution-infinite-or-finite>, դիտվել է 10.05.2019:

16 Kevin Rawlinson, "Microsoft's Bill Gates insists AI is a threat," *BBC*, January 29, 2015, <https://www.bbc.com/news/31047780>, դիտվել է 10.05.2019:

է միջուկային զենքի հետ և կարծում է, թե այն միևնույն ժամանակ հեռանկարային է և վտանգավոր¹⁷։ Իսկ Սթիվեն Հոփինգն ասում էր. «Արհեստական բանականության լիակատար զարգացումը կարող է հանգեցնել մարդկային ցեղի վերացմանը։ Հենց մարդիկ զարգացնեն արհեստական բանականությունը, այն կսկսի ինքնուրույն թռիչքաձև զարգանալ։ Մարդիկ, որոնք սահմանափակված են կենսաբանական դանդաղ էվոլյուցիայով, չեն կարող մրցակցել և դուրս կմղվեն»¹⁸։ Շվեդ փիլիսոփա, Օքսֆորդի համալսարանի պրոֆեսոր Նիք Բոստրոմն իր՝ «Սուպերբանականություն. ճանապարհները, վտանգները, ռազմավարությունները»¹⁹ գրքում պնդում է, որ եթե մեքենայական ուղեղները հիմնական բանականության հարցում գերազանցեն մարդկանց, ապա այս սուպերբանականությունը կարող է փոխարինել մարդկանց՝ որպես երկրի վրա գերիշխող կյանքի ձև։ Բանական մեքենաները կարող են ավելի արագ բարելավվել իրենց ընդունակությունները, քան մարդիկ, և դրա արդյունքը կարող է լինել մարդկության համար էքզիստենցիալ աղետը։

Ամերիկացի հայտնի պետական գործիչ, դիվանագետ Հենրի Քիսինջերը ևս կիսում է արհեստական բանականության զարգացման հետ կապված մտահոգությունները։ «Արհեստական բանականությունը որոշակի ընդունակությունների տիրապետում է ավելի արագ և ավելի հստակ, քան մարդիկ, և ժամանակի ընթացքում կարող է կրճատել մարդկային կարողությունը և մարդկային պայմանները՝ այն դարձնելով տվյալներ», - նշում է Քիսինջերը։ Ըստ նրա՝ մարդկային հասարակությունը պատրաստ չէ արհեստական բանականության աճին²⁰։ Հետաքրքրական է հատկապես Ռուսաստանի նախագահ Վլադիմիր Պուտինի մեկնաբանությունը. «Արհեստական բանականությունը ոչ միայն Ռուսաստանի, այլևս ողջ մարդկության ապագան է։ Այն գալիս է ոչ միայն իր հսկայական հնարավորություններով, այլև սպառնալիքներով, որոնք դժվար է կանխատեսել։ Ով կլինի այս ոլորտում առաջատարը, կդառնա աշխարհի ղեկավարը»²¹։

17 Jon Christian, "Bill Gates Compares Artificial Intelligence to Nuclear Weapons," March 19, 2019, <https://futurism.com/bill-gates-artificial-intelligence-nuclear-weapons>, դիտվել է 10.05.2019:

18 Cellan-Jones Rory, "Stephen Hawking warns artificial intelligence could end mankind," *BBC*, December 2, 2014, <https://www.bbc.com/news/technology-30290540>, դիտվել է 10.05.2019:

19 Nick Bostrom, *Superintelligence: Paths, Dangers, Strategies* (Oxford: Oxford University Press, 2014), 352.

20 Henry Kissinger, "How the Enlightenment Ends: Philosophically, intellectually - in every way - human society is unprepared for the rise of artificial intelligence," *The Atlantic*, June 2018, <https://www.theatlantic.com/magazine/archive/2018/06/henry-kissinger-ai-could-mean-the-end-of-human-history/559124/>, դիտվել է 10.05.2019:

21 Alexei Druzhinin, "'Whoever leads in AI will rule the world': Putin to Russian children on Knowledge Day," *Russia Today*, September 1, 2017, <https://www.rt.com/news/401731-ai-rule-world-putin/>, դիտվել է 12.05.2019:

Ինքնավար զենքերի վտանգը զանգվածային սպանությունների և ցեղասպանություն իրագործելու հարցում

Եթե գերբնական մեքենաներից եկող վտանգը ապագային վերաբերող հարց է, ապա ներկայումս ամենից մեծ վտանգը ներկայացնում են ինքնավար զենքերը, որոնք կարող են ձեռք բերել սեփական գիտակցություն: Ավելի մտահոգիչ է այն հանգամանքը, որ այն կարող է հայտնվել այնպիսի անհատի և կառավարության ձեռքում, որը մարդու կյանքը չի գնահատում²²: Եվ ամենավտանգավորն այն է, որ արհեստական բանականությունը կարող է կիրառվել ցեղասպանության պլանավորման համար²³:

2015 թ. ամռանը Արգենտինայում տեղի ունեցած արհեստական բանականության վերաբերյալ միջազգային համատեղ համաժողովին, որին մասնակցում էին Apple ընկերության համահիմնադիր Սթիվ Վոզնյակը, Tesla ընկերության հիմնադիր Իլոն Մասկը, գիտնական Ստիվեն Հոփինգը և ավելի քան 1.000 մասնագետներ, տեսնում էին պատին փակցված ցուցանակ, որի գրա գրված էր հետևյալը. «Եթե արհեստական բանականության տեխնոլոգիաները շարունակում են անսպասելիորեն զարգանալ՝ վերածվելով, ասենք, ինքնավար զենքերի համակարգերի, որոնք գործում են առանց մարդու մասնակցության, ի վերջո կիրականացնեն ցեղասպանություն և էթնիկ գտումներ»²⁴:

Կոնֆերանսի բացմանը՝ հուլիսի 28-ին, հրապարակվեց արհեստական բանականության և ռոբոտատեխնիկայի ոլորտում հետազոտություններ կատարող մասնագետների բաց նամակը: Նամակում նշվում էր, որ ինքնավար զենքերը դառնալու են վաղվա կալաշնիկովը: Ի տարբերություն միջուկային զենքի՝ ինքնավար զենքերը չեն պահանջում թանկ և դժվարամատչելի հումք, ու բոլոր կարևոր ռազմական ուժերի կողմից կդառնան էժան և համատարած արտադրության առարկա: Ոլորտի մասնագետները զգուշացնում էին, թե միայն ժամանակի հարց է, որ ինքնավար զենքերը հայտնվեն սև շուկայում, ահաբեկիչների ու այն բռնապետների ձեռքում, որոնք ցանկանում են ավելի խիստ վերահսկել իրենց ենթակա բնակչությանը, ինչպես նաև այն բարձրաստիճան գիտնականների մոտ, որոնք ցանկանում են էթնիկական գտումներ իրականացնել: Քանի որ ինքնավար զենքերը լայն հնարավորություններ են ստեղծում սպանությունների, բնակչությանը ենթարկեցնելու և որոշակի էթնիկ խմբերին ընտրողաբար ոչնչացնելու համար, մասնագետները կարևորում էին դրանց արգելումը²⁵:

22 Bernard Marr, “Is Artificial Intelligence Dangerous? 6 AI Risks Everyone Should Know About,” *Forbes*, November 19, 2018, <https://www.forbes.com/sites/bernardmarr/2018/11/19/is-artificial-intelligence-dangerous-6-ai-risks-everyone-should-know-about/#1bedd0212404>, դիտվել է 14.05.2019:

23 Andrew Majot and Roman Yampolskiy, “Diminishing Returns and Recursive Self Improving Artificial Intelligence,” in *The Technological Singularity: Managing the Journey (The Frontiers Collection)*, ed. Victor Callaghan, James Miller, Roman Yampolskiy, Stuart Armstrong (Berlin: Springer: 2017), 142.

24 Jonathan Benson, “AI technology to eventually streamline the type of genocide and ethnic cleansing carried out by Planned Parenthood, experts warn,” September 01, 2015, https://www.naturalnews.com/051010_artificial_intelligence_abortion_robots.html#, դիտվել է 12.05.2019:

25 Արհեստական բանականության և ռոբոտատեխնիկայի ոլորտում հետազոտություններ կատարող մասնագետների բաց նամակը տե՛ս “Autonomous Weapons: an Open Letter from AI & Robotics

Այս համատեքստում հատկանշական կարող ենք համարել նաև այն հանգամանքը, որ ժամանակի ընթացքում մեծանում է ինքնավար զենքերի համակարգերի՝ թիրախները խոցելու հնարավորությունը՝ արագությունը, ճշգրտությունը և դրանք գործում են ավելի համակարգված²⁶։ Ինքնավար զենքերը կարող են գործել առանց մարդկային անմիջական մասնակցության՝ դրանով բացառելով մարդկային հույզերի առաջացումը թիրախները խոցելու ժամանակ։ Սակայն հույզերի բացակայության հանգամանքը դժվար է գնահատել, քանի որ հույզերի ազդեցության ներքո մարդիկ կարող են և վայրագություններ կատարել, և մարդասիրական գործողություններ²⁷։

Ուշագրավ է, որ ներկայում ուղղակիորեն սկսվում է արհեստական բանականությամբ սպառազինությունների համաշխարհային մրցավազք, որի մասնակիցներն են ԱՄՆ-ը, Չինաստանը, Ռուսաստանը, Մեծ Բրիտանիան, Ֆրանսիան, Իսրայելը և Հարավային Կորեան²⁸։ Սակայն, գլխավորապես այն տեղի է ունենում ԱՄՆ-ի, Չինաստանի և Ռուսաստանի միջև²⁹։

2017 թ. օգոստոսին «Թեսլա» (Tesla) ընկերության հիմնադիր Իլոն Մասկի և արհեստական բանականությամբ զբաղվող բրիտանական DeepMind ընկերության հիմնադիր Մուստաֆա Սուլեյմանի գլխավորությամբ 26 երկրներից արհեստական բանականության և ռոբոտաշինական ընկերությունների հիմնադիր 116 մասնագետ բաց նամակ հղեց ՄԱԿ-ին՝ կոչ անելով կանխելու սպառազինությունների մրցավազքը մահացու ինքնավար զենքերի արտադրության ոլորտում։ Իրենց նամակում ընկերությունների հիմնադիրներն ընդգծում էին «Որոշակի սովորական սպառազինությունների մասին» կոնվենցիայի³⁰ վերանայման անհրաժեշտությունը։ Նրանք նշում էին, որ մահացու ինքնավար զենքերը կարող են դառնալ որպես «պատերազմի երրորդ հեղափոխություն»՝ վառողի հայտնագործումից և միջուկային զենքից հետո, ապա մատնանշում, որ դանդաղ գործելու ժամանակ չկա³¹։

Ըստ «Մայքրոսոֆթ» ընկերության նախագահ Սմիթի՝ մահացու ինքնավար զենքերի

Researchers,” <https://futureoflife.org/open-letter-autonomous-weapons/?cn-reloaded=1>, դիտվել է 13.05.2019:

26 Jayshree Pandya, “The Weaponization Of Artificial Intelligence,” *Forbes*, January 14, 2019, <https://www.forbes.com/sites/cognitiveworld/2019/01/14/the-weaponization-of-artificial-intelligence/#484da7953686>, դիտվել է 12.02.2020:

27 Paul Scharre, *Army of None: Autonomous Weapons and the Future of War* (New York, London: W. W. Norton & Company, 2018), 232.

28 Ավելի մանրամասն տե՛ս Kirsten Gronlund, “State of AI: Artificial Intelligence, the Military and Increasingly Autonomous Weapons,” *Future of Life Institute*, May 9, 2019, <https://futureoflife.org/2019/05/09/state-of-ai/>, դիտվել է 17.05.2019:

29 Ավելի մանրամասն տե՛ս Karla Lant, “China, Russia and the US Are in An Artificial Intelligence Arms Race,” September 12, 2017, <https://futurism.com/china-russia-and-the-us-are-in-an-artificial-intelligence-arms-race>, դիտվել է 17.05.2019:

30 “Convention on Prohibitions or Restrictions on the Use of Certain Conventional Weapons Which May Be Deemed to Be Excessively Injurious or to Have Indiscriminate Effects,” [https://www.unog.ch/80256EDD006B8954/\(httpAssets\)/51609D467F95DD5EC12571DE00602AED/\\$file/CONVENTION.pdf](https://www.unog.ch/80256EDD006B8954/(httpAssets)/51609D467F95DD5EC12571DE00602AED/$file/CONVENTION.pdf), դիտվել է 16.05.2019:

31 “An Open Letter to the United Nations Convention on Certain Conventional Weapons,” <https://www.cse.unsw.edu.au/~tw/ciair/open.pdf>, դիտվել է 18.05.2019:

օգտագործումը նոր խնդիրներ է առաջ բերում, որոնք կառավարությունների կողմից պետք է հրատապ դիտարկվեն³²:

Նշենք, որ պետությունների դիրքորոշումն այս հարցի շուրջ միանշանակ չէ: Օրինակ՝ 2015 թ. Մեծ Բրիտանիայի կառավարությունը դեմ է հանդես եկել ինքնավար զենքերի արգելքին՝ նշելով, թե միջազգային մարդասիրական իրավունքն արդեն իսկ ապահովում է բավարար կարգավորում այդ ոլորտի համար: Միացյալ Թագավորությունն ընդգծել է, որ չի զարգացնում մահացու ինքնավար զենքերի համակարգեր, իսկ Մեծ Բրիտանիայի զինված ուժերի կողմից կիրառվող բոլոր զենքերը մշտապես գտնվում են մարդկային վերահսկողության ներքո³³:

2016 թ. Չինաստանը հրապարակում է ծրագրային մի փաստաթուղթ, որում նշվում է ինքնավար զենքերի խնդրի անհամապատասխանությունը միջազգային իրավունքի համընդհանուր նորմերին՝ մատնանշելով այս հարցի ուղղությամբ միջազգային իրավունքում փոփոխությունների անհրաժեշտությունը, սակայն՝ ոչ կիրառման արգելումը³⁴: Այս մոտեցումն է ցուցաբերում նաև Ռուսաստանը՝ չընդունելով ինքնավար մահացու զենքերի օգտագործման արգելքը³⁵: 2019 թ. մարտի 29-ին ՄԱԿ-ի՝ ինքնավար մահացու զենքերի արգելման խնդրին նվիրված Ժնևյան նիստին Ավստրալիան, Իսրայելը, Ռուսաստանը, Մեծ Բրիտանիան և ԱՄՆ-ը դեմ քվեարկեցին արգելքին³⁶:

2019 թ. աշնանը ՄԱԿ-ի Գլխավոր ասամբլեայի 74-րդ նստաշրջանում առնվազն 41 երկիր բարձրացրել է մարդասպան ռոբոտների արգելքի հարցը, իսկ Ավստրիան, Չինաստանը, Կուբան, Էկվադորը և Գվատեմալան նորից հանդես են եկել մահացու ռոբոտների կիրառումն արգելող միջազգային նոր պայմանագրի տեքստի կազմման անհրաժեշտության իրենց կոչերով³⁷:

2017 թ. Հարվարդի համալսարանի Բելֆեր կենտրոնն իր զեկույցում նշում

32 Wesley Hudson, "Killer robots 'unstoppable' rise will put public safety 'at risk' claims Microsoft chief," *Daily Express and Sunday Express*, September 22, 2019, https://www.express.co.uk/news/uk/1181082/killer-robots-latest-technology-news-microsoft-ai-lethal-autonomous-weapons?fbclid=IwAR2nnpEYHljlqgbvYK_bstXU08Dk1YpPMRYnuZDD-nBH46Trls-J5-9qXo8, դիտվել է 05.02.2020:

33 Owen Bowcott, "UK opposes international ban on developing 'killer robots,'" *The Guardian*, April 13, 2015, <https://www.theguardian.com/politics/2015/apr/13/uk-opposes-international-ban-on-developing-killer-robots>, դիտվել է 05.02.2020:

34 "The position paper submitted by the Chinese delegation to CCW 5th Review Conference," [https://www.unog.ch/80256EDD006B8954/\(httpAssets\)/DD1551E60648CEBBC125808A005954FA/%24file/China's+Position+Paper.pdf](https://www.unog.ch/80256EDD006B8954/(httpAssets)/DD1551E60648CEBBC125808A005954FA/%24file/China's+Position+Paper.pdf), դիտվել է 19.05.2019:

35 E&T editorial staff, "Russia rejects potential UN 'killer robots' ban, official statement says," *E&T magazine*, December 1, 2017, <https://eandt.theiet.org/content/articles/2017/12/russia-rejects-potential-un-killer-robots-ban-official-statement-says/>, նաև՝ "Group of Governmental Experts of the High Contracting Parties to the Convention on Prohibitions or Restrictions on the Use of Certain Conventional Weapons Which May Be Deemed to Be Excessively Injurious or to Have Indiscriminate Effects," November 10, 2017, <https://admin.govexec.com/media/russia.pdf>, դիտվել է 19.05.2019:

36 Gayle Damien, "UK, US and Russia among those opposing killer robot ban," *The Guardian*, March 29, 2019, <https://www.theguardian.com/science/2019/mar/29/uk-us-russia-opposing-killer-robot-ban-un-ai>, դիտվել է 08.11.2019:

37 "High-level concerns on killer robots at UN," October 30 2019, <https://www.stopkillerrobots.org/2019/10/unga74/>, դիտվել է 08.11.2019:

է, որ գրեթե անհնար է կանխել արհեստական բանականության ռազմական նպատակներով կիրառման ընդլայնումը³⁸: Եթե անգամ ՄԱԿ-ը կանխարգելիչ արգելք սահմանի ինքնավար զենքերի տեխնոլոգիաների կիրառման վրա, իսկ այդ արգելքը պատշաճորեն պահպանվի պետությունների կողմից, այդուհանդերձ, հարցը չենք կարող փակված համարել:

Որոշակի խմբի անդամների թիրախավորելու նպատակով դրոնների կիրառման հարցը

Ինքնավար զենքի համակարգի միջոցով հնարավոր է թիրախավորել ռասայական խմբի անդամներին՝ նրանց ոչնչացնելու համար: Օրինակ՝ դրոնը, օգտագործելով նկարներ ճանաչող ծրագրեր, կարող է տարբերակել անհատի էթնիկ պատկանելիությունը (գույնով, ֆիզիկական առանձնահատկություններով):

Հատկանշական է, որ այսպիսի դրոններ ստեղծելու փորձերը կարող են արվել մասնավոր հաստատություններում՝ սահմանափակ բյուջեով և կարճ ժամանակահատվածում³⁹: Ազգությամբ ճապոնացի ամերիկացի գիտնական, տեսական ֆիզիկայի ոլորտի փորձագետ Միցիո Կակուն զգուշացնում է ռազմականացված դրոնների համակարգերի բացարձակ վտանգավորության մասին. ըստ նրա՝ մոտ ապագայում դրոնը կճանաչի մարդու կառուցվածքը, որը հնարավորություն կընձեռնի սպանելու թիրախավորվածին⁴⁰:

Այսօր արդեն լայնորեն գործածվում են դեմքերի նույնականացման համակարգերը: Ավելին՝ Չինաստանն արդեն փորձարկում է վերահսկողական այնպիսի համակարգեր, որոնք թույլ են տալիս հայտնաբերել նույնիսկ մարդկանց հույզերը⁴¹: Ռուսական «Էն Թեչլաբ» (NtechLab) ընկերությունն օգտագործում է դեմքը ճանաչելու համակարգը՝ մարդկանց էթնիկ պատկանելությունը որոշելու նպատակով⁴²:

«Future of Life Institute»-ը⁴³ 2017 թ. թողարկեց «Slaughterbots» անվանումով

38 Greg Allen, Taniel Chan, *Artificial Intelligence and National Security*, Belfer Center for Science and International Affairs (Cambridge: Belfer Center for Science and International Affairs Harvard Kennedy School, 2017), 3, <https://www.belfercenter.org/sites/default/files/files/publication/AI%20NatSec%20-%20final.pdf>, դիտվել է 22.05.2019:

39 Amir Husain, *The Sentient Machine: The Coming Age of Artificial Intelligence* (New York, London, Toronto, Sydney and New Delhi: Simon and Schuster, 2018), 106-107.

40 Jolene Creighton, “Physicist Michio Kaku Has Some Powerful Predictions for the Future,” March 26, 2018, <https://futurism.com/michio-kaku-prominent-futurist-predictions>, դիտվել է 19.05.2019:

41 Adrian Potoroaca, “China is testing ‘emotion recognition’ systems in Xinjiang,” *TechSpot*, November 4, 2019, <https://www.techspot.com/news/82611-china-testing-emotion-recognition-systems-xinjiang.html?fbclid=IwAR23iKberEyXCPTMHGGSD6Ss-dhNHun2yJCkm5KPBU5ilvIhuF4jvZ3yKkE>, դիտվել է 19.01.2020:

42 Ryan Daws, “Russian startup is building a controversial ‘ethnicity-detecting’ AI,” *AI News*, May 25, 2018, <https://artificialintelligence-news.com/2018/05/25/russian-ethnicity-detecting-ai/>, դիտվել է 19.05.2019:

43 Future of Life Institute-ը (թարգմ.՝ Կյանքի ինստիտուտի ապագան) Բոստոնի տարածքում գործող կամավոր հետազոտական և կրթական կազմակերպություն է, որն աշխատում է մեղմել մարդկության առջև ծառացած էկզիստենցիալ վտանգները՝ մասնավորապես առաջադեմ արհեստական բանականության էկզիստենցիալ ռիսկը:

ավելի քան 7 թույլ տևողություն ունեցող մի տեսանյութ, որում գիտաֆանտաստիկ ձևով ցույց է տրվում, թե ինչպես է տեխնոլոգիական ընկերությունը ներկայացնում ամենաառաջադեմ ինքնավար դրոնները, որոնք փոքր խաղալիքի չափ ունեն և զինված են պայթուցիկով: Դրոնները նույնականացնում են թիրախներին և սպանում՝ նախապես ծրագրավորված չափանիշների և համակարգված տվյալների հիման վրա⁴⁴: «Future of Life Institute»-ի նպատակն էր ցույց տալ ինքնավար զենքերի կիրառումից եկող վտանգները: Եվ ամենակարևորն այն է, որ խոսքը գնում է ոչ միայն վաղվա հիպոթետիկ վտանգների մասին, այլև հենց ներկայի:

Ուշագրավ է, որ Չինաստանն արդեն սկսել է Մերձավոր Արևելքում վաճառել ինքնավար մահացու ռոբոտներ: Օրինակ՝ չինական «Չիյան» (Ziyan) ընկերությունը միջազգային գնորդներին ակտիվորեն գովազդում է «Բլոուֆիշ Ա3» (Blowfish A3) ուղղաթիռը, որն ինքնաձիգով զինված ինքնավար դրոն է:

ԱՄՆ պաշտպանության նախարար Մարկ Էսպերն այս առնչությամբ նշել է, թե չինական զենք արտադրողները վաճառում են անօդաչու այնպիսի սարքեր, որոնք լիարժեք ինքնավարություն ունեն, ներառյալ՝ նպատակային մահացու հարվածներ հասցնելու հնարավորություն⁴⁵:

Յեղասպանական քարոզչության համար սոցիալական մեդիայի և կեղծ տեսանյութերի կիրառման խնդիրը

Էթնիկ և կրոնական փոքրամասնությունների նկատմամբ արհեստական բանականության միջոցով ցեղասպանություն իրագործելու քարոզչության առումով լուրջ գործիք կարող է ծառայել նաև սոցիալական մեդիան: Սոցիալական ցանցերում քարոզչություն իրականացնելով ալգորիթմների և անձնական տվյալների օգտագործման արդյունքում նույնականացված անհատներին արհեստական բանականությունը կարող է թիրախավորել և տարածել ցանկալի տեղեկատվություն⁴⁶: Արհեստական բանականությունը կարող է ապակայունացնել համացանցային նորությունների ողջ էկոհամակարգը, իսկ սոցցանցերը կարող են դառնալ առաջին զոհերը⁴⁷:

Այս համատեքստում հատկապես ուշագրավ է այն հանգամանքը, որ արդեն գրանցվել են դեպքեր, երբ սոցիալական ցանցերը օգտագործվել են ցեղասպանական նպատակներով, ինչպես 2016-2017 թթ. Մյանմարում զինծառայողների կողմից ֆեյս-

44 Future of Life Institute, “Slaughterbots,” November 13, 2017, https://www.youtube.com/watch?v=HipTO_7mUOW, դիտվել է 15.09.2019:

45 Patrick Tucker, “SecDef: China Is Exporting Killer Robots to the Mideast,” *Defense One*, November 5, 2019, <https://www.defenseone.com/technology/2019/11/secdef-china-exporting-killer-robots-mideast/161100/?oref=DefenseOneFB&fbclid=IwAR30HzhnTAAzJH-LUggx4leuTuR71PdLEcunjhzinNnwnxDVplJJ7OH EA7s>, դիտվել է 29.09.2019:

46 Bernard Marr, “Is Artificial Intelligence Dangerous? 6 AI Risks Everyone Should Know About,” *Forbes*, November 19, 2018, <https://www.forbes.com/sites/bernardmarr/2018/11/19/is-artificial-intelligence-dangerous-6-ai-risks-everyone-should-know-about/#1bedd0212404>, դիտվել է 19.01.2020:

47 Victor Tangerman, Elon Musk, “‘Advanced AI’ Will Manipulate Social Media,” 26 September 2019, <https://futurism.com/the-byte/elon-musk-ai-manipulate-social-media>, դիտվել է 05.11.2019:

բուքյան հարթակը կիրառվեց որպես էթնիկ գտումների գործիք, և կեղծ էջերի միջոցով մահմեդական ռոհինջանների նկատմամբ հակաքարոզչություն իրականացվեց⁴⁸:

«Մայքրոսոֆթ»-ը 2016 թ. ներկայացրեց արհեստական բանականությամբ օժտված Թայ անունով չաթբոտը⁴⁹, որը պետք է մեկնաբանություններ աներ սոցիալական մեդիայում և սովորեր, թե ինչպես կարող է շփվել ուրիշների հետ: Սակայն երկար ժամանակ պահանջվեց, որ «Մայքրոսոֆթ»-ը Թայի ազգայնամոլ քարոզչական գրառումները ջնջի. չաթբոտը վերջնականապես արգելափակվեց, երբ սկսեց ցեղասպանությունը քարոզող գրառումներ անել⁵⁰:

Արհեստական բանականության գործիքները կարող են լայնորեն օգտագործվել նաև «դիպֆեյք» (deepfake) անունը ստացած կեղծ տեսանյութերի տարածման համար: Օրինակ՝ տեսագործիքների կիրառմամբ անհատի դեմքը կարող է տեղադրվել մեկ այլ մարդու մարմնի վրա⁵¹, արհեստական բանականության միջոցով հնարավոր է ընդօրինակել նաև մարդու ձայնը⁵²:

Այսպիսի տեսանյութերը դժվարությամբ հայտնաբերվող առույիտ կամ վիդեո թվային մանիպուլյացիաներ են, որոնցում կարող են պատկերվել մարդիկ, և ներկայացվել, թե նրանք կատարում են այսպիսի արարքներ կամ այսպիսի խոսքեր են արտաբերում, երբ նրանք նման նման խոսքեր չեն ասել կամ արարքներ չեն կատարել: Այսպիսի բովանդակությամբ կեղծ տեսանյութերը կարող են օգտագործվել որպես գործիք ահաբեկչական խմբավորումների կողմից սադրիչ գործողությունների իրականացման և թիրախային լսարանների հույզերի վրա ազդելու նպատակով՝ անգամ ցույց տալով, թե իբր իրենց հակառակորդները ցեղասպանական գործողություններ են իրականացնում: Օրինակ՝ «Իսլամական պետություն» ահաբեկչական խմբավորման կողմից արհեստական բանականության միջոցով կարող է ստեղծվել կեղծ տեսանյութ, որում պատկերված կլինեն ԱՄՆ զինծառայողներ, ովքեր կրակում են խաղաղ բնակիչների ուղղությամբ կամ քննարկում են մզկիթը ռմբակոծելու պլանը, որի միջոցով ահաբեկչական խմբավորումը պայմաններ կստեղծի, որ մարդիկ թշնամաբար

48 Paul Mozur, "A Genocide Incited on Facebook, With Posts From Myanmar's Military," *New York Times*, October 15, 2018, <https://www.nytimes.com/2018/10/15/technology/myanmar-facebook-genocide.html>, դիտվել է 08.03.2020:

49 Չաթբոտը համակարգչային ծրագիր է, որն իրականացնում է հաղորդակցություն տեքստային կամ ձայնագրված նամակների միջոցով: Չաթբոտերը սովորաբար օգտագործվում են տարբեր գործնական երկխոսություններում, ինչպես, օրինակ՝ հաճախորդների սպասարկումը կամ տեղեկատվություն ստանալը:

50 James MacDonald, "Microsoft is deleting its AI chatbot's incredibly racist tweets," *Business Insider*, April 15, 2016, <https://daily.jstor.org/when-the-ai-promotes-genocide/>, դիտվել է 23.05.2019: Տե՛ս նաև Rob Price, "Microsoft is deleting its AI chatbot's incredibly racist tweets," *Business Insider*, March 24, 2016, <https://www.businessinsider.com/microsoft-deletes-racist-genocidal-tweets-from-ai-chatbot-tay-2016-3?r=UK&IR=T>, դիտվել է 28.05.2019:

51 Kevin Roose, "Here Come the Fake Videos, Too," *New York Times*, March 4, 2018, <https://www.nytimes.com/2018/03/04/technology/fake-videos-deepfakes.html>, դիտվել է 28.05.2019:

52 Bahar Gholipour, "New AI Tech Can Mimic Any Voice," *Scientific American*, May 2, 2017, <https://www.scientificamerican.com/article/new-ai-tech-can-mimic-any-voice/>, դիտվել է 29.05.2019:

տրամադրվեն ԱՄՆ-ի նկատմամբ, ինչպես նաև համալրեն իրենց շարքերը⁵³:

Մեծ վտանգ կա, որ արհեստական բանականության գործիքների միջոցով կեղծ տեսանկյունների տարածումը կկիրառվի հենց ատելություն քարոզելու նպատակով: Եվ պատահական չէ, որ որոշ մասնագետներ մտավախություն ունեն, որ կեղծ տեսանկյունները կարող են օգտագործվել ինչպես ապատեղեկատվություն տարածելու, այնպես էլ բռնություն, պատերազմ կամ ցեղասպանություն հրահրելու նպատակով⁵⁴: Եվ մասնավորապես կրոնական կամ էթնիկ փոքրամասնությունների առաջնորդների մասնակցությամբ կեղծ տեսանկյունները հնարավորություն կտան այդ փոքրամասնության նկատմամբ ատելության մթնոլորտ ստեղծել և ցեղասպանություն հրահրել⁵⁵:

Վերոնշյալները դեռևս չեն ներկայացնում բոլոր այն վտանգները, որոնք կարող են առաջանալ արհեստական բանականության կիրառության արդյունքում. արհեստական բանականության առաջընթացն աներևակայելի արագ է⁵⁶ և կարծում ենք՝ դրա հետագա զարգացումն իր հետ կարող է նոր վտանգներ բերել:

Արհեստական բանականության էթիկայի մշակման անհրաժեշտությունը

Կանխատեսելով արհեստական բանականության առաջ բերած հնարավոր վտանգները՝ դեռևս 1950 թ. գրող Այզեկ Ազիմովն առաջ է քաշել ռոբոտատեխնիկայի երեք օրենք. ռոբոտը չպետք է վնասի մարդուն, ռոբոտը պետք է կատարի մարդկանց հրահանգած հանձնարարությունները՝ բացառությամբ առաջին օրենքի և ռոբոտը պետք է պաշտպանի իր գոյությունն այնքանով, որքանով այն չի հակասի առաջին և երկրորդ օրենքներին⁵⁷:

Գեղարվեստական ստեղծագործություններում, որտեղ ռոբոտները պատասխանատվություն են ստանձնել ամբողջ մոլորակի և մարդկային քաղաքակրթությունների կառավարման համար, Ազիմովը ավելացրել է նաև չորրորդ կամ զրոյական (zereth) օրենքը՝ ռոբոտին արգելվում է վնասել մարդկությանը: Ռոբոտն իրավունք ունի վնասել անհատին միայն այն դեպքում, երբ դրանով կկանխի ողջ մարդկության գոյատևմանը

53 Robert Chesney and Danielle Citron, “Deepfakes and the New Disinformation War: The Coming Age of Post-Truth Geopolitics,” *Foreign Affairs*, January/February 2019, <https://www.foreignaffairs.com/articles/world/2018-12-11/deepfakes-and-new-disinformation-war>, դիտվել է 28.01.2020:

54 Council on Foreign Relations, “The Threat of Deep Fakes,” May 31, 2019, <https://www.youtube.com/watch?v=zhf2X1RyIWE>, դիտվել է 20.09.2019:

55 John Loeffler, “AI Deepfakes Are Now Easier to Make Than Ever, Pointing to Ominous Future,” *Interesting Engineering*, June 11, 2019, <https://interestingengineering.com/ai-deepfakes-are-now-easier-to-make-than-ever-pointing-to-ominous-future>, դիտվել է 28.01.2020:

56 Stephen Yuen, “Elon Musk on Artificial Intelligence: ‘The pace of progress in artificial intelligence is incredibly fast,’” November 18, 2014, <https://www.thesavvytechs.com/2014/11/18/elon-musk-artificial-intelligence-pace-progress-artificial-intelligence-incredibly-fast/>, դիտվել է 28.05.2019:

57 Isaac Asimov, “*Runaround.*” *I, Robot* (New York City: Doubleday, 1950), 40.

սպառնացող վտանգը⁵⁸։ Այդուհանդերձ, ժամանակակից ռոբոտատեխնիկայի մասնագետներն Ազիմովի վերոնշյալ օրենքների կիրառումն իրատեսական չեն համարում՝ հիմնական խնդրի երկիմաստության պատճառով⁵⁹։

Հնարավոր բացասական հետևանքները կանխելու նպատակով ոլորտի մասնագետներն անհրաժեշտ են համարում արհեստական բանականության էթիկայի մշակումը։ Այն ներառում է մեքենայական էթիկան և ռոբոտէթիկան, որը վերաբերում է մարդկանց «բարոյական վարքին»՝ արհեստական բանականությամբ օժտված էակների նախագծում, կառուցում, օգտագործում։

2016 թ. Ստանդարտների բրիտանական ինստիտուտը ռոբոտներ ստեղծողների համար հրապարակեց ռոբոտների էթիկական վարքագծի ապահովման վերաբերյալ «BS8611. ռոբոտներ և ռոբոտական սարքեր» անվանումով փաստաթուղթ, որում նշվում էր, որ «ռոբոտները չպետք է նախագծված լինեն մարդուն սպանելու կամ վնասելու համար, և այդ հարցում ոչ թե ռոբոտները, այլ մարդիկ են պատասխանատու։ պետք է հնարավոր լինի պարզել, թե ով է պատասխանատու ցանկացած ռոբոտի և դրա պահվածքի համար»⁶⁰։

Այսպիսով՝ մարդկության համար հրատապ է դարձել վերոնշյալ հնարավոր վտանգները կանխելու նպատակով օրենսդրական համապատասխան մեխանիզմների ստեղծումը՝ հիմնված մարդու իրավունքների պաշտպանության և էթիկայի չափորոշիչների վրա։ Սա նաև հնարավորություն կտա կանխարգելել ցեղասպանության և զանգվածային սպանությունների իրագործման հնարավոր վտանգները։

Թեև ցեղասպանագետները հետազոտություններ են իրականացրել տեղեկատվության և հաղորդակցման տեխնոլոգիաների դերի ուղղությամբ⁶¹, նույնիսկ առաջարկ կա ցեղասպանության կանխարգելման համատեքստում օգտագործել արբանյակների հնարավորությունները՝ մարդու իրավունքների խախտումների շուրջ հեռակա մշտադիտարկում անելու և վաղ նախագգուշացման համակարգ գործարկելու նպատակով⁶², սակայն արհեստական բանականության թեմային դեռևս համապատասխան անդրադարձ չի կատարվել։ Հետևաբար, հարկ է մտածել արհեստական բանականության համապատասխան ալգորիթմների մշակման ուղղությամբ, որոնք թույլ կտան հայտնաբերել և կանխարգելել ցեղասպանական միջավայրի ձևավորումը վաղ փուլում։

58 Keith Abney, "Robotics, Ethical Theory, and Metaethics: A Guide for the Preplexed," in *Robot Ethics: The Ethical and Social Implications of Robotics*, ed. Patrick Lin, Keith Abney, George A. Beke (Cambridge, Massachusetts, London: MIT Press, 2012), 43.

59 McCauley Lee, "AI Armageddon and the Three Laws of Robotics," *Ethics and Information Technology* 9, no.2 (2007): 153-164.

60 Amit Ray, *Compassionate Artificial Intelligence: Frameworks and Algorithms*, Compassionate AI Lab (An Imprint of Inner Light Publishers, 2018), 46-47.

61 Չանգվածային սպանությունների կատարման գործում տեղեկատվության և հաղորդակցման տեխնոլոգիաների դերակատարության մասին տե՛ս "Information and Communications Technologies in Mass Atrocities Research and Response," *Genocide Studies and Prevention* 11, no. 1, (2017): 128.

62 Andrew Marx, "Using Satellites to Detect Mass Human Rights Violations: A Call for International Community to Implement an Early Warning Detection System," in *Last Lectures on the Prevention and Intervention of Genocide*, ed. Samuel Totten (London and New York: Routledge, 2018), 171-179.

Narek M. Poghosyan

PhD in History

THE DEVELOPMENT OF ARTIFICIAL INTELLIGENCE AND RISKS FOR THE IMPLEMENTATION OF GENOCIDE AND MASS KILLINGS

SUMMARY

Key words: Artificial intelligence, lethal autonomous weapons, arms race, robots, genocide, social media, fake videos, artificial intelligence ethics.

This article illustrates the fact that the accelerating pace of technological development in recent years, in particular the use of artificial intelligence (AI) has raised serious concerns among experts that it, along with its advantages, can pose many dangers to mankind when machines think and make decisions for them. Particularly risky is that the development of AI and autonomous weapons can be used to target certain national, religious and racial groups, thereby increasing the risk of genocide and mass killings.

In the present day, when the global arms race using AI is a reality, many technologists are turning their attention to banning lethal autonomous weapons. In addition, in order to prevent the potential negative consequences of AI development, industry professionals have suggested the need to develop an AI code of ethics.

The creation of appropriate legislative mechanisms based on ethics and the protection of human rights has already become an urgent matter for mankind to prevent the possible risks posed by the development of robotics and artificial intelligence technology.

Since artificial intelligence is used to detect and prevent various types of crime, it is therefore necessary to use the opportunities provided by it to prevent genocide and mass murder.

Нарек М. Погосян

кандидат исторических наук

РАЗВИТИЕ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА И РИСКИ ДЛЯ ОСУЩЕСТВЛЕНИЯ ГЕНОЦИДА И МАССОВЫХ УБИЙСТВ

РЕЗЮМЕ

Ключевые слова: искусственный интеллект, смертоносное автономное оружие, гонка вооружений, роботы, геноцид, дроны, социальные сети, фальшивые видео, этика искусственного интеллекта.

В статье показано, что ускоряющиеся темпы развития технологий в последние годы, в частности использование искусственного интеллекта, вызывают серьезные

опасения у экспертов, что он, наряду с его преимуществами, может создать много опасностей для человечества, когда машины думают и принимают решения для себя. Особенно рискованным является тот факт, что разработка искусственного интеллекта и автономного оружия может использоваться для нацеливания на определенные национальные, религиозные, расовые группы, что повышает риск геноцида и массовых убийств.

В настоящее время, когда глобальная гонка вооружений с использованием искусственного интеллекта становится реальностью, многие технологи обращают свое внимание на запрет смертоносного автономного оружия. Кроме того, чтобы предотвратить возможные негативные последствия развития искусственного интеллекта, профессионалы отрасли предложили разработать этику искусственного интеллекта.

Создание соответствующих законодательных механизмов, основанных на защите прав человека и этики, уже стало насущной задачей для человечества, чтобы предотвратить возможные риски, связанные с развитием робототехники и технологий искусственного интеллекта.

Поскольку искусственный интеллект используется для обнаружения и предотвращения различных видов преступлений, поэтому необходимо использовать возможности, предоставляемые искусственным интеллектом для предотвращения геноцида и массовых убийств.