

www.eai.or.kr

EAI 스페셜리포트

AI와 신문명 표준: 군사도전 ①

인공지능-핵무기 넥서스(AI-Nuclear Nexus)와 세계군사질서 전망

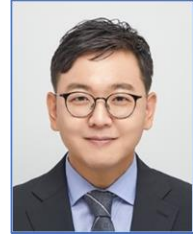
김양규(EAI; 서울대)

AI와 신문명 표준: 군사도전 ①

인공지능-핵무기 넥서스(AI-Nuclear Nexus)와 세계군사질서 전망

김양규

EAI 수석연구원, 서울대 정치외교학부 강사



I. 서론

2024년 서울에서 3차례 주요한 인공지능(Artificial Intelligence: AI) 규범에 관한 국제회의가 개최됨에 따라 한국 사회 내 AI에 대한 관심이 높다. 지난 3월에는 “AI·디지털 기술과 민주주의”를 주제로 제3차 <민주주의 정상회의(Summit for Democracy)>가, 5월에는 제2차 <인공지능 서울 정상회의(AI Seoul Summit)>가 열린데 이어, 오는 9월 9-10일에는 제2차 <인공지능의 책임 있는 군사적 이용에 관한 고위급회의(Summit on Responsible Artificial Intelligence in the Military Domain: REAIM)>가 개최된다. 첨단기술 관련 규범을 논의하는 국제회의를 한국에서 연이어 개최한다는 것은 국제사회 내 한국의 위상과 영향력을 제고할 수 있는 좋은 기회임에 틀림없다. 그러나, 한국이 지속적으로 글로벌 중추국가로서 국제사회 내 리더십을 발휘하기 위해서는 무분별한 첨단기술의 군사적 활용을 통제하는 보편적 규범을 제시할 뿐 아니라, 대한민국 영토 전체를 사정범위에 두고 선제핵무기 타격을 위협하고 있는 북한의 실존적 위협에 대처하는 고민을 동시에 해나가야 한다.

본 스페셜리포트는 AI의 군사적 이용 문제 가운데 북핵위협으로 인해 한국에게 가장 중요한 영향을 미칠 핵무기와 AI가 접합되는 AI-핵무기 넥서스(AI-Nuclear Nexus)를 살펴본다. 첫째, AI가 무엇을 의미하는 것인지, 이것을 군사영역에 적용할 경우 어떠한 변화가 일어나는 것인지 개괄적으로 논의한다. 둘째, AI 기술이 핵무기 역량 및 전략과 결합될 경우 이는 기존의 공격-방어 균형(Offense-Defense Balance)에 어떤 영향을 미치게 되는지 살펴본다. 셋째, 이러한 분석을 토대로 AI-핵무기 넥서스가 향후 세계 군사질서를 어떻게 변화시키게 될 것인지 전망한다.

II. AI의 개념 및 군사적 이용: 작전 속도 증가의 승수효과

AI가 무엇인가에 관해 공식적으로 합의된 개념은 없으나 대부분의 연구들은 특정 ‘과제’를 수행함에 있어 상황 파악, 패턴 식별, 결론 도출, 예측, 계획, 학습, 의사소통 등과 같은 ‘인간 지능’이 요구되는 일을 할 수 있는 기계(machine)라는 점을 공통적으로 지적한다(Horowitz 2018, 40; Haenlein and Kaplan 2019, 5; US Department of Defense 2019, 5; Congressional Research Service 2020, 2). 록펠러 재단(Rockefeller Foundation)의 지원 하에 1956년 다트머스대학교 AI 하계 연구 프로젝트(Dartmouth Summer Research Project on Artificial Intelligence: DSRPAI)를 시작으로 AI 연구는 “봄(AI Spring)”과 “여름(AI Summer)”을 맞이했지만, 1980년대 초 미국, 영국 등 주요국 정부들이 “경험 많은 아마추어”의 수준을 넘지 못하는 AI 성능에 실망하여 재정지원을 끊으면서 1990년대까지 겨울(AI Winter)을 맞이했다. IBM이 개발한 딥블루(Deep Blue)가 1997년 세계 체스 챔피언 카스파로프(Garry Kasparov)를 꺾으면서 잠시 주목을 받았으나, 특정 규칙을 기반으로 하는 전문가 시스템(Expert System) 방식으로는 AI가 특정한 영역을 넘어서 다양한 분야에서 성능을 발휘하는 데까지 이르지 못했기에 그 한계가 뚜렷했다(Haenlein and Kaplan 2019, 6-8).

2015년 구글이 개발한 알파고(AlphaGo)가 체스보다 훨씬 복잡하다고 여겨지는 바둑에서 커제와 이세돌을 차례로 꺾고 세계 바둑을 제패하면서 AI 연구는 두번째 여름을 맞았다(Haenlein and Kaplan 2019, 9). 입력한 값을 바탕으로 결과를 추론하는 과정에서 파라미터(parameter)라고 불리는 수많은 가중치(weight)와 편향(bias)을 적용하여 기계 스스로 답을 찾게 만드는 인공신경망(Artificial Neural Network: ANN) 방식의 머신러닝(machine learning)이 새로운 돌파구를 마련한 것이었다(Bode et al. 2024, 3-5). 특히, 2010년대 중반 이후 SNS의 활성화로 인해 접근 가능한 데이터가 폭발적으로 늘어났고, 데이터를 분산하여 병렬 계산할 수 있기에 머신러닝에 적합한 GPU 칩(예. A100, H100, B100)이 발전함에 따라 컴퓨터 연산 능력이 비약적 향상되었다.

이를 토대로 기존에는 이론 수준에 머물러 있던 AI 모델들을 실제 현실에서 구현해낼 수 있는 조건이 마련되었다. 이제는 어떠한 형태의 미가공 데이터(raw data)를 입력하더라도 기계 스스로 데이터를 분류하고 학습의 길을 찾는 딥러닝(Deep Learning)이 가능해졌고, OpenAI의 ChatGPT와 같이 ‘X’를 해달라고 지시(prompt)만 하면 이야기, 음악, 그림, 동영상, 코딩, 전략까지 순식간에 만들어 내는 생성형 AI(Generative AI: GAI)의 시대를 맞이했다. 생성형 AI의 등장으로 인해 인공지능은 폭넓은 영역에 걸쳐 기존에는 불가능했던 것들을 “가능케 해주는 기술(enabler)”이라는

점이 분명해졌다. 따라서 AI는 탱크, 잠수함 같은 군사 기술이라기 보다는 전기, 내연기관 같은 범용 기술이다(Horowitz 2018, 41). 이런 맥락에서 AI 기술은 기존 지휘통제체제(Command, Control, Communication, computer and Intelligence: C4I) 및 무기체계와 결합하여 “전력증강자(force multiplier)” 역할을 할 수 있다(Johnson 2019, 150).

범용 기술인 생성형 AI의 발전은 군사안보 차원에서는 어떤 변화를 가져오게 될까? 관련 연구들이 공통적으로 지적하는 것은 ‘속도’와 관련되어 있다. 앞서 살펴본 바와 같이, 머신러닝과 딥러닝에 기초한 AI는 엄청난 양의 정보를 순식간에 분석하여, 이를 토대로 주어진 상황에서 선택할 수 있는 옵션 가운데 가장 효과적인 방안을 찾아내는데 가장 뛰어난 면을 보인다. 이는 21세기 들어 꾸준히 진행된 군사 혁신이 (1) 속도, (2) 거리, (3) 정확도의 문제였다는 점(Metz 2000, 73-81; Lieber and Press 2017; Schneider and Macdonald 2024, 174-177)을 고려할 때 AI가 ‘속도’에 있어 혁명적인 변화를 가져온다는 것은 동 기술이 군사적으로 활용될 경우 군사 혁신 경쟁에서 우위에 설 수 있음을 의미한다.

특히, 빠른 속도로 전장 정보를 처리하게 되면, 표적 탐지, 식별, 추적은 물론 상대방의 전술적 움직임에 대한 대응속도까지 빨라지게 되므로, 이는 ‘거리’와 ‘정확도’ 측면에서도 승수효과(multiplier)를 가져올 수 있다. 나아가 최근 미국의 “통합억지(integrated deterrence)”나 중국의 “지능화전(intelligentized warfare)” 개념들이 강조하는 육·해·공, 우주, 사이버, 비군사 등 국가 안보역량의 통합 및 다영역(multi-domain) 작전 수행 차원에서도 AI 기술을 접목하면 가장 효과적인 방식의 조합 계산을 신속히 할 수 있기 때문에 매우 중요해진다(김양규 2023; 2024). 이는 앞으로 ‘어떤 국가가 다른 국가들보다 먼저 AI 기반 군사력을 보유 하느냐’에 따라 해당 국가 군사작전의 성패가 결정되는 상황을 맞이할 가능성이 크다는 점을 시사한다.

현재 AI 분야에서 독보적인 위치를 점하고 있는 미국이 AI의 군사적 활용 계획을 살펴보면 이러한 측면이 그대로 드러난다. 미 국방부(U.S. Department of Defense)가 2023년 6월 발표한 『데이터, 분석, 인공지능 활용 전략서(Data, Analytics, and Artificial Intelligence Adoption Strategy)』의 부제는 “의사결정 속도 증진의 이점(Accelerating Decision Advantage)”으로 되어 있다. 본 전략서의 서두에서는 미군이 AI에 주목해야 하는 이유를 “더 좋은 결정을 더 빠르게 해주기(enable leaders to make better decisions faster)” 때문이라고 강조한다(U.S. Department of Defense 2023, 3). AI 기술이 군사작전의 속도에 큰 영향을 준다는 점을 보여주는 대목이다.

아울러, (1) 전장 환경 인식 및 이해(battlespace awareness and understanding), (2) 적응력 있는 부대 계획 및 적용(adaptive force planning and application), (3) 빠르고, 정확하고, 회복력 있는 킬체인(fast, precise, and resilient kill chains), (4) 회복력 있는 군수지원(resilient sustainment support), (5) 효율적 작전임무 조직(efficient enterprise business operations)이 갖춰져야 미국이 추구하는 “의사결정에서의 우위(decision advantage)”를 달성할 수 있다고 강조하면서, 이를 위해서도 AI 기반 인간-기계 협업과 신속한 정보분석 처리가 꼭 필요하다고 역설한다.

구체적으로 『데이터, 분석, 인공지능 활용 전략서』는 미국 AI 군사화의 전략 목표로 (1) ‘양질의 데이터 구축’을 바탕으로 한 (2) “합동전쟁수행(joint warfighting) 능력”을 신장하여 (3) ‘책임 있는 AI 의 군사적 이용’을 달성하는 것을 제시한다. 여기서 합동전쟁수행은 “전술에서 전략 수준에 이르기까지 합동 능력의 공백을 메우는 것(address Joint capability gaps at the operational to strategic levels)”으로 조직간 상호운용성(interoperability)을 높임으로써 전쟁수행능력을 강화하는 것을 의미한다. 이를 종합적으로 고려하면, 의사결정 속도를 높일 경우, 전쟁 수행 능력, 작전 지속성 및 성공을 위한 기초 역량을 갖추게 될 뿐 아니라, 합동전쟁수행 능력 달성을 위한 조건을 마련할 수 있다.

이런 측면에서 AI 를 군사적으로 이용할 경우 그 함의는 매우 포괄적일 가능성이 높다. AI 를 구체적으로 어떻게 군사적인 측면에서 활용할 수 있는가에 관해 기존 연구는 (1) AI 를 전략적으로 사용할 것인가, 전술적으로 사용할 것인가, (2) 인간이 AI 를 감독할 것인가, 기계에게 자율성을 줄 것인가를 기준으로 하여 아래 표와 같이 AI 의 군사적 사용 형태를 구분한다(Lushenko 2023).

〈표 1〉 AI 의 군사적 이용 유형

	전술적(Tactical) 활용	전략적(Strategic) 활용
인간 감독 (Human Oversight)	켄타우로스 전쟁 (Centaur Warfighting)	모자이크 전쟁 (Mosaic Warfare)
기계 자율 (Machine Oversight)	미노타우로스 전쟁 (Minotaur Warfare)	AI 지휘관 (AI-general)

AI를 ‘전술적’으로 사용한다는 것은 센서(sensors)를 통해 취득한 대량의 정보를 빠르게 처리하여 표적에 대한 대응을 매우 빠르게 진행함으로써 전장에서 “탐지 및 정밀타격(sensor-to-shooter)”에 걸리는 시간을 최소화하는데 중점을 두는 것을 말한다. 한편, AI를 ‘전략적’으로 사용한다는 것은 군사적 목표 달성을 위해 어떤 방식의 전쟁을 치르고, 어떤 전력을 조합한 뒤 투사할 것인지에 대한 옵션을 파악하기 위해 AI를 활용한다는 것이다. 통제 수준은 기계가 자율적으로 판단하고 실행하도록 허용할 것인지, 아니면 기계는 인간의 판단에 도움이 될 수 있는 분석결과를 제공할 뿐 최종 판단은 인간이 하는 방식으로 운영할 것인지에 따라 구분한 것이다. 이러한 기준에 따라 분류하면 총 네 가지 형태의 AI 군사적 이용 모델이 가능하다.

‘켄타우로스’는 그리스-로마 신화에 등장하는 괴인으로 상체는 인간, 하체는 말의 형태를 띠고 있다. 즉, AI를 전술적으로 사용하여 전장에서 군사작전의 효율성을 담보할 수 있는 선택지를 마련하되, 최종 결정은 인간이 하는 형태이다. ‘미노타우로스’는 머리는 괴물에 몸은 사람 형태를 가진 신화적인 괴물로 순찰병력의 배치에서부터 전투기 편대 배치까지 전장에서 작전 수행을 위한 판단을 모두 기계가 하고 인간은 이를 그대로 따르는 방식이다. AI가 드론을 통제하는 자율무기체계(Autonomous Weapon System: AWS)가 여기에 해당한다. ‘모자이크 전쟁’은 적의 움직임을 예측하고 취약점을 최대한 공략할 수 있는 아군 전력의 최적화된 조합 파악을 위해 AI를 전략적 수준에서 활용하지만 최종적인 결정은 인간이 하는 형태를 일컫는다. 앞서 언급한 현재 미국이 추진 중인 AI의 군사적 활용 계획은 루센코(Paul Lushenko)의 분류에 따르면 모자이크 전쟁에 속한다. 끝으로 ‘AI 지휘관’은 군사력 활용에 대한 국가의 전반적인 전략 결정을 모두 AI에 위탁하고 인간이 개입하지 않는 형태이다.

루센코는 미군 현장 지휘관들을 대상으로 앞의 네 가지 형태의 AI 활용 모델에 대한 선호도 조사를 실시했다. 응답자들은 “신뢰도” 측면에서는 모자이크 전쟁(1 순위), 켄타우로스 전쟁(2 순위), 미노타우로스 전쟁(3 순위), AI 지휘관(4 순위) 순의 선호도를 보였는데, “실제 어떤 형태의 AI 모델을 군이 채택하기를 바라는가?”에 대한 질문에 대해서는 미노타우로스 전쟁(1 순위), 모자이크 전쟁(2 순위), 켄타우로스 전쟁(3 순위), AI 지휘관(4 순위) 순의 지지도를 보였다. 이러한 결과는 AI 기술이 군사영역에 적용되었을 때 얼마나 안정성을 보일지가 아직 불확실한 상황에서 보편적인 인간의 심리는 전술 및 전략 레벨에서 모두 인간이 AI를 통제하는 방식을 선호한다는 경향성을 드러낸다.

흥미로운 것은 구체적으로 군사무대에서 AI가 어떤 역할을 해주기를 바라는지에 대한 문제에 있어서는 개별 전장에서 전쟁을 수행하는 ‘전술’ 차원의 문제는 AI에게 판단을 맡겨 자율무기체계

활용을 통한 효율성 증대를, 국가안보 ‘전략’ 차원에서 전쟁 계획을 세우고 국가 안보 역량을 어떻게 통합할 것인가의 문제는 인간이 주도권을 가지고 AI 는 인간의 판단을 돕는 보조적 역할에 머무는 것을 더 선호하는 다소 상충적인(ambivalent) 경향을 보인다는 점이다. 이는 대부분의 군 지휘관들이 AI 기술의 안정성에 대해 완전한 확신을 가지고 있지는 않지만, 최소한 전술적 수준에서 AI 를 활용하는 것에 대해서는 큰 거부감을 가지고 있지 않음을 시사한다. 물론 본 조사는 미군 지휘관들만을 대상으로 하였고 어떤 기준에서 응답자를 선발하였는지 알 수 없기 때문에 이를 성급하게 일반화하기에는 무리가 있다. 그럼에도 불구하고, 만약 실제로 AI 가 제공하는 의사결정 과정에서 작전 속도 증대와 효율성 강화를 군 지휘관들이 경험하게 되면, 군사 무대에서 AI 의 역할은 점점 더 확대될 가능성이 크다.

다음 절에서는 AI 의 군사적 이용이 현재까지 인류가 개발한 무기체계 가운데 가장 혁명적인 것이라고 여겨지는 핵무기의 효율성 및 향후 핵 전략에 어떤 영향을 미치게 될지 살펴본다. 이 때, AI 기술의 도입이 “공격-방어 균형 문제(Offense-Defense Balance) (Jervis 1978)”에서 어떤 영향을 미쳐 핵 전략 영역에서 공격과 방어 중 무엇이 우위를 가져올 가능성이 더 높을지 살펴보고, 궁극적으로는 AI 의 군사적 이용이 핵무기와 같은 기존 무기를 보완(complement)하는지, 혹은 대체(replace)하는지의 문제를 함께 검토한다.

III. AI와 핵무기:

대가치타격(Countervalue) 중심에서 대군사타격(Counterforce) 중심으로 변화

1. 핵무기 혁명(Nuclear Revolution)과 2차 공격능력의 중요성

핵무기는 2 차대전 당시 히로시마와 나가사키에서 사용된 이후 인류 역사상 두 번 다시 사용되지 않았을 정도로 그 사용이 정당화되기 어려운 ‘지나치게 강력한(overkill)’ 무기체계이다. 핵무기는 투하되는 즉시 발생하는 태양 복사 에너지 수준의 강력한 열, 급속히 공기를 가열함으로 인해 발생하는 극단적인 대기압의 변화와 그로 인해 발생하는 태풍같은 강력한 바람, 그리고 이어서 떨어지는 방사능 낙진으로 인해 탑재된 핵탄두의 위력에 따라 투하 지점 반경 수십 킬로미터에 이르기까지 생명체의 생존을 불가능하게 만드는 끔찍한 무기이다(Wolfson and Dalnoki-Veress 2022). 전투원과 비전투원을

구분하지 않는 가공할만한 살상력을 매우 빠른 속도로 적에게 투사할 수 있기 때문에(Fetter 1991; Pape 1996), 핵무기를 사용한 전쟁에서 승자는 있을 수 없으며, 무슨 수단을 동원해서라도 핵전쟁을 피하는 것이 국가안보 전략의 최우선 순위가 된다. 그런데 역설적으로 바로 이러한 측면 때문에 핵무기 시대 국가들은 모두 전쟁을 피하고 우발적 확전을 경계하는 신중한 안보정책을 추진하였다. 이것은 무려 40 년 동안이나 미소간 직접적인 군사적 충돌이 없었던 냉전기 전략적 안정성의 역사에서 보듯, 인류에게 혁명적인 변화를 가져다 주었다(Carnesale et al. 1983; Jervis 1989; Waltz 2003).

이제까지 핵무기에 대한 안보전략을 구축함에 있어 가장 중요한 개념은 “2 차 공격능력(second-strike capability)”이었다(Wohlstetter 1959). 이는 상대방으로부터 핵 공격을 받은 이후에도 핵무기로 반격하여 상대방에게 보복할 수 있는 능력을 의미하는데, 쌍방이 모두 2차 공격능력을 보유하여 핵무기 사용시 공멸에 이르게 되는 상태를 “상호확증파괴(Mutual Assured Destruction: MAD)”라고 부른다. MAD 가 형성되면 대립하는 두 핵 국가가 서로에 대하여 동일하게 ‘용납할 수 없는 수준의 피해(unacceptable damage)’를 입힐 수 있기 때문에 누구도 먼저 핵으로 상대방을 위협할 수 없게 된다. 서로가 상대방 국민의 안녕과 생존을 불모로 잡아 역설적으로 양국 모두 자국의 안보를 지킬 수 있는 상황이 되는 것이다(Jervis 1989). 눈에는 눈, 이에는 이, 자국 국민 생명 위협에는 타국 국민 생명 위협으로 맞선다고 하여 2 차 공격능력을 “대가치타격(countervalue attack)”능력이라고도 부른다(Kahn 1960).

한편, 상대방의 2 차 공격능력을 무효화할 수 있는 능력을 1 차 공격능력(first-strike capability)이라고 하는데, 이는 상대방이 가진 핵전력을 파괴하는 것에 초점을 맞추고 있어 “대군사타격(counterforce attack)”이라고도 한다. 미소 냉전시기 핵전략 전개 역사를 보면, 핵을 보유한 국가들이 서로에 대한 2 차 공격능력의 강화를 추진하면 전략적 안정성이 높아지고, 1 차 공격능력 확보를 위해 노력할 경우 전략적 불안정성이 커지는 양상을 보였다(Kaplan 1983; Jervis 1989; Freedman 2003). 대표적인 것이 미사일 방어체계(Missile Defense: MD)인데, MD 가 겉으로 보기에는 자국 영토와 국민을 적국의 핵 공격으로부터 보호하는 방어적인 수단인 것처럼 보이지만, 100%의 신뢰성을 가진 MD 를 구축하는데 성공할 경우 이는 상대방의 2 차 공격능력을 무효화 시키기 때문에, MAD 가 붕괴되는 효과를 가져온다. 이것이 미소간 최초의 핵군축 시도였던 전략 무기 제한 협상(Strategic Arms Limitation Talks: SALT) 과정에서 탄도탄 요격미사일 조약(Anti-Ballistic Missile Treaty: ABM Treaty)이 함께 논의되고, 동 조약이 SALT I 과 함께 1972 년 최초의 미소 군축 합의가 된 이유이다.

2. AI의 군사적 이용과 핵무기: 1차 공격능력 강화와 핵 비대칭성 심화

따라서 AI 기술의 도입이 핵전략 차원에서 어떤 변화를 가져오는지 판단하기 위해서는 동 기술이 1차 공격능력(대군사타격능력)과 2차 공격능력(대가치타격능력) 중 어디에 더 큰 이점을 가져오는지 살펴볼 필요가 있다. 핵 전략 차원에서 1차 공격능력은 ‘공격 우위’를 강화하는데 기여하는 변수이고, 2차 공격능력은 ‘방어 우위’를 담보하는 요소이다.

본격적인 논의에 앞서 하나의 중요 원칙으로 견지해야 할 점은 기술 변수 하나만으로 공격-방어 균형의 변화를 설명할 수 없다는 점이다. 전술한 바와 같이 AI를 군사적으로 이용할 때 가질 수 있는 가장 큰 장점은 작전 ‘속도(tempo)’의 증가에 있다. AI는 표적 식별, 작전 환경 분석 및 이해, 표적 타격을 위한 최적의 무기체계 조합 산출 등의 영역에서 인간의 인지 능력을 압도적으로 능가하는 효율을 낼 수 있기 때문에 정확하면서도 빠른 결정을 내릴 수 있고, 그 효과는 군사력 사용의 전 주기에 미칠 가능성이 높다. 그러나, 이 능력 자체는 공격에 기여할 수도, 방어에 기여할 수도 있다. 따라서 AI 변수 하나만으로 공격 우위 또는 방어 우위의 시대가 열린다고 단정짓기는 어렵다. 결국 중요한 것은 기술 자체가 아니라 그 기술을 어떻게 운용 하는가의 문제이기 때문이다(Biddle 2023).

이를 염두에 두고, “인공지능-핵무기 넥서스(AI-Nuclear Nexus)”가 공격-방어 균형 차원에서 미래 전장에 어떤 변화를 가져올지 살펴보자. 먼저 공격 혹은 대군사타격 능력에 AI가 기여할 수 있는 부분은 정보·감시·정찰(Intelligence, Surveillance, and Reconnaissance: ISR), 핵 지휘통제(Nuclear Command and Control: NC2), 그리고 재래식 전력 기반 대군사타격 작전(conventional counterforce operations)이다(Johnson 2023, 18-23; 78-84; 87-90).

첫째, AI로 인해 강화된 ISR 능력은 산개, 이동, 엄폐, 보호 등의 방법으로 핵무기 생존성 및 2차 공격능력을 담보하려는 핵 국가에게 큰 위협이 된다. 중국, 북한, 심지어 러시아의 경우에도 3대 핵전력(Nuclear Triad) 가운데 잠수함 역량에는 한계가 많다. 따라서 이들 국가들은 대륙간탄도미사일(Intercontinental Ballistic Missile: ICBM)을 강화된 미사일 격납고(hardened missile silos)나 지하시설(Underground Facilities: UGF)에 두어 은폐 또는 방호를 시도하거나, 이동식발사대(Transporter Erector Launchers: TELs)를 활용하여 적 탐지에서 벗어나려고 하는 등의 방법으로 자국 핵 자산의 생존력을 높이려 한다(Lieber and Press 2017). 그런데 이러한 전술은 미국이 드론, 인공위성 등 기존 ISR 자산을 AI 기술에 접목할 경우 정찰 능력이

극대화되기에 쉽게 무력화될 수 있다. 한 예로, 탐지 및 추적이 어려워 이제까지 궁극적인 2 차 공격능력 확보 수단이라고 여겨져 왔던 핵잠수함조차, 미국의 방위고등연구계획국(Defense Advanced Research Projects Agency: DARPA)에서 잠수함을 식별하고 추적하는 씨헌터(Sea Hunter) 무인 드론을 개발함에 따라 그 취약성이 높아지고 있다. 즉, 식별과 탐지 역량의 제고는 곧 정밀 타격 능력 강화로 이어지기 때문에 정보·감시·정찰은 1 차 공격능력 강화에 기여하는 기술이다.

둘째, 핵 지휘통제(NC2)의 경우 AI 로 강화된 사이버 및 전자기전 공격에 취약해질 가능성이 있다. 예를 들어, AI 기술력에서 우위에 있는 국가가 적국에게 지능형 지속 공격(Advanced Persistent Threat: APT) 작전을 감행할 경우, 표적이 된 국가의 지휘통제 시스템에 대해 지속적인 컴퓨터 해킹 프로세스들이 가동된다. 이 경우 적의 사이버 방호 체계의 약점을 찾아내고, 탐지 불가능한 바이러스, 멀웨어 등을 심어 오작동을 유도할 수 있다. 가짜 이미지나 표적 정보를 심어 데이터 포이즈닝(data poisoning)이나 스푸핑(spoofing) 전략을 구사할 수도 있다. 이런 사이버 공격의 경우 재래식 무기를 동원한 정밀 타격 능력과 달리 상대방 지휘체계에 대한 물리적인 파괴를 시도하지는 않지만, 궁극적으로는 NC2 가 제대로 기능하지 못하기 만들기 때문에 실질적인 효과는 물리적 공격과 차이가 없다. 따라서 사이버 및 전자기전 능력 강화 역시 1 차 공격능력 강화로 보아야 한다. 특히, 사이버 방호 측면에서 열세에 있다고 생각하는 국가들은 경보 즉시 발사(Launch on Warning: LOW) doktrinen을 채택하여 이러한 문제를 해결하려 들 수 있어 불안정성이 매우 커진다.

셋째, 재래식 전력 기반 대군사타격은 전형적인 1 차 공격능력에 속한다. AI 기술의 도입은 무인 무기체계의 정확성을 비약적으로 향상시킬 수 있어 적 방어력 돌파에 용이하다. 특히, 중국의 반접근 지역거부(Anti-Access Area Denial: A2/AD) 전략과 같이 기존의 유인 시스템으로 돌파하기에는 너무나 많은 비용을 치러야 했던 적 방어체계에 대해서도 매우 효과적인 대항 수단이 된다. 역으로 미사일 방어체계에 AI 기반의 자동표적식별(Automatic Target Recognition: ATR)이 추가되면 탐지·추적·요격 능력이 급격하게 강화되고, 여기에 드론 스웜밍(drone swarming)을 결합시키면 거부에 의한 억지(deterrence by denial) 능력을 크게 높인다. 전술한 바와 같이 이러한 능력 향상은 자국의 전략 자산 및 영토를 보호하는 방어적인 조치처럼 보이지만, 실제로는 상대방의 2 차 공격능력을 무효화하는 효과가 있기 때문에 공격적인 속성을 가진다.

그렇다면 AI 기술의 도입은 핵전략 차원에서 공격의 절대 우위를 가져와 불안정성만 커지게 만드는 것인가? 반드시 그렇다고 볼 수 없다. AI 기술은 핵무기 전략 차원에서 방어적 능력을 강화하는 부분도 존재한다. 예를 들어, AI 기술은 장거리 정찰과 복잡한 지형에 대한 실시간 정보

수집 및 분석 능력을 비약적으로 성장시키기에, 방어하는 측의 조기경보 정확도를 포함하여 상황인식 및 대비 능력 강화에 기여한다. 적 방어체계를 무력화하는데 사용될 수 있는 감시정착 능력은 방어하는 국면에서는 적의 기습 공격을 무력화시킬 수도 있는 것이다.

앞서 언급한 핵 지휘통제 능력에 대한 사이버전에 있어서도 자국 사이버 방어체계의 취약점을 발견하고, 적의 데이터 포이즈닝과 스푸핑 시도에 대한 식별 등을 AI 기술의 접목시킨 사이버 방어능력 신장을 통해 대응할 수 있다. 드론 스위밍의 경우에도 일반적으로 방어하는 측이 자국 군사 자산에 대한 정보를 공격하는 측보다 훨씬 많이 가지고 있기 때문에, 이 데이터의 비대칭성은 딥러닝에 의한 AI 성능 측면에서 방어가 공격보다 더 유리한 결과로 이어질 수 밖에 없다. 아울러, 현재 드론의 제한적인 장거리 작전 수행 능력, 기동 속도 등을 고려할 때 AI 기술이 접목된 드론 역량은 공격하는 측보다 방어하는 측을 더 유리하게 만든다는 연구도 있다(King 2024). 따라서 AI 기반 핵 자산 방호 능력은 AI-재래식무기 대군사타격은 작전에 대한 훌륭한 대항 수단을 제공한다.

뿐만 아니라, 러시아의 “죽음의 손(Dead hand)” 발사 코드와 같이, 기존의 핵 억지 전략의 신뢰성 구축에서 가장 큰 맹점으로 지목되어 온 핵무기를 동원한 보복 ‘의지(resolve)’의 근본적 한계를 AI로 극복할 수 있는 길이 열렸다. 기존 핵 연구에서 반복적으로 지적되었던 핵 억지(nuclear deterrence)의 문제는 방어국이 물리적으로 2차 공격능력을 가지고 있다 하더라도 핵 보복 자체의 비합리성으로 인해 MAD의 안정성에 의구심을 가질 수 밖에 없다는 점이었다. “(핵을 사용한) 세계 전쟁에서 패배하는 것보다 더 끔찍한 일은 핵 전쟁에서 살아남는 것이다(the only thing worse than losing a global war was winning one)”라는 아이젠하워 대통령의 유명한 말처럼, 이미 핵 공격에 의한 피해를 입은 국가 입장에서는 상대방을 동일하게 핵으로 공격하여 얻을 수 있는 효용이 거의 없기 때문이다.

이 문제를 해결하기 위해 쉘링(Thomas Schelling)은 게임 모델에 입각하여 핵 통제력을 의도적으로 낮추는 “Threats That Leave Something to Chance” 방식의 대응을 제안하기도 했고, 최근 연구(McDermott et al 2017)는 핵 보복이 신뢰성을 가지기 위해서는 상대방이 우리에게 가한 고통을 응징함으로써 얻을 수 있는 ‘심리적 만족감’에 기대어 핵 보복 의지를 시그널링 하는 방안을 제시하기도 한다. 그런데 “죽음의 손”처럼, ‘적이 핵으로 공격’ 하였는데 ‘핵 통제권을 가진 주체와 연결이 끊어진 상황’이라면 ‘자동적으로 핵무기로 보복한다’는 방식으로 프로그래밍하여 AI에게 핵 보복 능력을 통제할 권리를 부여하게 되면, 핵 보복과 2차 공격능력의 신뢰성 문제를 완전히 극복할 수도 있다.

따라서, ‘AI-핵무기 넥서스’는 핵 전략에 있어 반드시 공격 우위를 가져오지도, 방어 우위를 가져오지도 않는다. 어떤 모델을 어떻게 적용하느냐에 따라 전혀 다른 결과를 가져올 수 있다.

3. AI의 군사적 이용과 핵무기의 미래: 핵무기 효용성 감소

추가로 논의해야 할 중요한 이슈는 AI 기술의 도입이 실제 전장에서 핵무기를 사용할 필요성을 급격히 낮추게 되는 문제이다. AI 기술은 신속대응과 정밀타격 능력을 신장시키는 효과가 가장 강하고, 이 경우 핵무기의 과도한 살상 능력은 전술적으로 큰 의미가 없다. 예를 들어, AI 기반 대가치타격 능력은 민간인에 대한 대량 살상 보다는 적 지휘시설에 대한 정밀 타격 및 지휘관 제거에 더 적합하다. 2022년 발표된 미국의 핵 태세검토보고서(Nuclear Posture Review: NPR)에서 “김정은 정권의 종말”을 이야기 한 것도 이와 연관된 것으로 보인다.

AI 기반 대군사타격 능력은 또한, 적의 취약한 부분에 대한 정밀 타격으로 상대방의 2차 공격능력을 무력화하거나, 그 무력화를 시도하는 적 재래식 전력에 대한 대응 및 파괴 등에서 높은 효율을 발휘한다. 이렇게 되면, 핵무기는 전략적 수 싸움에서 그 맥락을 형성하는 배경적인 변수로만 작동할 뿐, 실제 전장에서 핵무기를 사용해야 할 필요성을 느끼는 상황 자체가 발생하지 않을 가능성이 크다.

동시에 앞서 살펴본 AI-핵무기 넥서스의 공격-방어 우위 계산에서 보듯, 쌍방이 비슷한 수준의 AI 능력을 보유한 경우, 이것이 공격하는 측과 방어하는 측 누구에게도 일방적인 우위를 부여하지는 않는다. 그러나 만약 미국과 북한처럼 한 쪽(미국)의 AI 능력이 상대방(북한)보다 압도적으로 우위에 있을 경우, 우세한 측은 자국의 2차 공격능력에 대한 방어력은 극대화하면서 동시에 상대방의 핵 자산도 효과적으로 파괴할 수 있기 때문에 압도적인 수준의 1차 공격능력을 가지게 될 가능성이 크다. 이 경우 북미간 무력 충돌 및 강압정책의 대립 상황에서 북한의 핵무기는 미국의 AI-핵무기 통합 군사력 앞에서 어떠한 효용도 갖지 못하게 될 가능성이 크다.

IV. AI-핵무기 넥서스와 미래 군사질서 전망: 미중전략경쟁의 파멸적 귀결 가능성

종합적으로 고려할 때, AI 기술의 군사적 사용은 기존 ‘핵무기 능력’을 증폭시키거나 핵무기의 중요성을 강화하는 전력 증강자(multiplier)로 기능하기 보다는, 기존 ‘재래식 전력’의 효율성을 급격히 신장시켜 핵 사용의 필요성 자체를 원천적으로 제거하는 방향으로 작동할 가능성이 커 보인다. 이런 맥락에서, AI 기반 군사력은 핵무기를 강화하는 보완재이기 보다는 핵무기의 전략적 필요성과 효용성을 크게 감소시키는 대체재인 측면이 더 크다.

AI 의 성능이 데이터의 양과 질에 영향을 받는다는 점, 그리고 AI 가 제대로 기능하기 위해서는 머신러닝에 적합한 연산능력이 반드시 필요하다는 점을 고려한다면 장기적으로 AI 능력 차원에서 미중간 비대칭성은 커질 것이다. 실제 전장 상황에서만 수집할 수 있는 이미지, 영상, 장비 성능에 관한 데이터(Horowitz 2018, 52-54) 차원에서 미국이 타의 추종을 불허하는 데이터를 보유하고 있다는 점, 그리고 미국이 세계 첨단반도체 생산 장비 공급망의 90 퍼센트를 주도하고 있다는 점에서(Allen 2023) 기존 초강대국인 미국이 최첨단 AI 능력 차원에서도 독점적 지위를 유지할 가능성이 매우 크다. 이렇게 되면, 미국 주도의 패권질서는 계속되고 세계 정치질서의 다극화는 일어나지 않을 것이다. 물론, 중국이 연산능력이나 군사용 데이터 차원에서 미국이 걸어온 경로와 전혀 다른 새로운 발전 방향을 찾게 되면(예. 양자컴퓨팅, 중국산 드론의 판매 및 보안정보 탈취) 미국이 중국과 경쟁에서 압도적 우위를 점하지 못하는 미래도 배제할 수는 없다.

그러나 이러한 장기적 미국 주도의 세계군사질서는 중단기적으로 상당히 불안정한 시기를 지나게 될 가능성이 매우 높다. AI 의 군사적 이용과 관련하여 현재 기술적 차원에서 대표적 한계로 AI 가 인간과 같은 수준의 유추나 적응을 하지 못하는 것에서 비롯되는 “불안정성(brittleness)”의 문제(Johnson 2023, 12-14)와 인간에 대한 폭력 사용의 결정을 기계가 하게 하는 ‘윤리적 문제’(Bode et al. 2024, 8-9)가 있다. 그런데 이것 외에도 미중 AI 경쟁은 그 과정에서 군사적 불안정성이 크게 높아지는 두가지 동학에서 자유롭기 힘든 심각한 문제가 있다.

첫째, ‘핵 얽힘(nuclear entanglement)’ 문제이다. 이는 현재 미중 핵경쟁과 관련하여 중국, 러시아를 비롯한 여러 나라들이 핵탄두와 비핵탄두 모두 장착 가능한 이중용도(dual-use) 투발수단과 핵무기 운용 부대와 재래식 전력 운용 부대를 조직적으로 결합하는 등 핵무기의 배치 및 운용 단계에서 재래식 전력과 핵 전력을 의도적으로 얽히게 만들어 제한적 핵 탄두수의 효과를 극대화하려 할 때 지적된 문제이다. 핵 얽힘은 실제로는 적국이 자국의 재래식 전력만을 공격할

의도로 군사작전을 전개하더라도 그와 얽혀 있는 핵무기까지 취약해지기 때문에 해당 국가는 보유한 핵 전력을 ‘사용하거나 상실하거나(use-it-or-lose-it)’의 상황에 놓이게 된다. 이는 우발적 핵 전쟁 가능성을 크게 높인다(Acton et al 2017; Talmadge 2017). 재래식 전력을 사용한 제한전이 쉽게 핵전쟁으로 비화되는 것이다.

비슷한 문제가 AI-핵 넥서스 구축 과정에서도 발생할 수 있다. 충분한 데이터가 공급되지 않으면 AI 는 현재 핵 얽힘 전략을 구축한 상대국이 실전배치한 미사일, 폭격기, 잠수함 자산들이 재래식 탄두를 탑재하고 있는지, 핵탄두를 탑재하고 있는지 정확히 판별해 내기 어렵다. 특히 AI 머신러닝 기술이 기반을 둔 파라미터 방식의 추론이 가지고 있는 블랙박스 문제(Johnson 2023, 17-18)로 인해 AI 는 정보 분석만 돕고 모든 전술적 판단은 인간이 내리도록 엄밀하게 통제하고 있는 상황이라고 하더라도 핵 얽힘에 따른 불확실성이 AI 로 인해 증폭되어 더욱 심각한 불안정성을 야기할 수 있다. 즉, 인공지능이 왜 특정 군사적 행동을 추천하는 판단에 이르렀는지 추가 설명을 제공하지 않는 경우, 혹은 제공하더라도 그것이 참인지 교차검증하는 것이 불가능할 경우 제한된 시간 내에 어떤 결정을 내려야 하는지를 놓고 인간 지휘관은 상당한 심리적 압박을 받을 수 밖에 없다. 이 경우 결국은 AI 의 판단을 따르는 쪽으로 선택할 가능성이 높다.

둘째, 핵 얽힘 전장 환경의 불확실성 속 AI 기술의 블랙박스 문제가 결합될 때 ‘의도치 않은 확전’ 가능성이 크게 높아진다. 존슨(James Johnson)은 미중이 모두 AI 기반 방위태세를 갖추고 있는 상황에서 미국 고위 관료가 대만을 방문한 것에 대해 중국이 무력시위를 할 경우 대만해협에서 미중 간 사이버 작전 수준의 상호 견제 행동이 어떻게 핵전쟁으로 비화될 수 있는지 실감나는 시나리오를 제공한다(Johnson 2023, 1-3). 미중 전략경쟁이 전쟁 문턱을 넘어서는 대응을 하는 핵심적인 기점은 AI 기반 감시정찰 및 전략적 판단 시스템이 ‘조기 확전 우세(early escalation dominance)’를 위해 선제공격을 추천하는 것이다. 특히, 미국 보다 핵탄두 수가 적어 핵 얽힘 전략을 채택하고 있는 중국 입장에서는, 자국 AI 기반 국방시스템이 현재 미국이 중국의 대미 핵 억지 역량 무력화를 위해 중국의 재래식 전력 뿐 아니라 핵무기 또한 노리고 있기 때문에 조속히 선제 공격을 감행해야 한다고 추천하면 이를 무시하기가 상당히 어렵다. 그리고 중국의 이러한 상황을 잘 알고 있는 미국 입장에서도 중국의 다음 대응이 우주에 있는 미국 인공위성 자산에 대한 공격이나 괌 기지 내 미국 핵심 자산에 대한 초음속미사일 공격이라고 자국의 AI 국방시스템이 예측하면, 이에 대응한 대군사타격 선제공격을 감행할 가능성이 높다. 이는 모두 쉽게 인도-태평양 지역에서 전쟁이 발발할 수 있는 여건을 제공한다.

V. 결론

따라서 AI-핵 넥서스 경쟁에서 장기적으로 미국이 우위를 점할 가능성이 매우 크지만, 중단기적으로는 핵 얽힘 및 의도치 않은 확전에 의해 미중 간에 심각한 무력 충돌이 발생할 가능성이 있다. 중국 입장에서 미국을 장기적으로 따라잡는 것이 불가능해지다고 느낄수록 더욱 더 현 시점에서 전쟁을 선택하는 것이 낫다는 압박을 받게 될 것이다. 지금부터 이러한 가능성에 대해 제대로 대비를 하지 않는다면, 좁게는 인도-태평양 지역, 넓게는 인류가 공멸에 이르는 경우의 수도 배제하기 어렵다. 이것이 바로 AI의 군사적 이용에 대한 보편 규범 수립이 절실한 이유이다. '핵무기 사용에 관한 결정에 AI를 관여시킬 것인지'에 관한 문제를 비롯하여, 파멸적 경로가 예측되는 '핵 얽힘'과 '의도치 않은 확전'에 관한 논의 등 미중이 다른 이슈들에 비해 합의를 볼 가능성이 큰 어젠다를 중심으로 AI 전략 대화를 진행할 필요가 있다. ■

참고문헌

- Acton, James, et al. 2017. *Entanglement: Chinese and Russian Perspective on Non-nuclear Weapons and Nuclear Risks*. Washington, DC: Carnegie Endowment for International Peace.
- Alarcon, Nefi. 2020. "OpenAI Presents GPT-3, a 175 Billion Parameters Language Model." July 7. <https://developer.nvidia.com/blog/openai-presents-gpt-3-a-175-billion-parameters-language-model/>. (검색일: 2023. 11. 20.)
- Allen, Gregory C. 2023. "Blocking China's Access to AI Chips Matters to U.S. National Security." CSIS Commentary. July 31. <https://www.csis.org/analysis/blocking-chinas-access-ai-chips-matters-us-national-security>. (검색일: 2024. 1. 4.)
- Bai, Yunyi. 2023. "Wang-Sullivan talks last more than 12 hours; Taiwan question takes up longest time, Chinese FM official told GT." *Global Times*. September 18. <https://www.globaltimes.cn/page/202309/1298392.shtml>. (검색일: 2023. 11. 20.)
- Biddle, Stephen. 2023. "Back in the Trenches: Why New Technology Hasn't Revolutionized Warfare in Ukraine." *Foreign Affairs*. August 10. <https://www.foreignaffairs.com/ukraine/back-trenches-technology-warfare> (검색일: 2024. 1. 4.)
- Bode, Ingvild, et al. 2024. "Algorithmic Warfare: Taking Stock of a Research Programme." *Global Society* 38(1): 1-23.
- Bremmer, Ian and Mustafa Suleyman. 2023. "The AI Power Paradox: Can States Learn to Govern Artificial Intelligence - before It's Too Late?" *Foreign Affairs* 102: 26.
- Carnesale, Albert, et al. 1983. *Living with Nuclear Weapons*. New York: Bantam Books.
- Congressional Research Service. 2020. "Artificial Intelligence and National Security." November 10. <https://sgp.fas.org/crs/natsec/R45178.pdf>. (검색일: 2024. 1. 16.)
- Elimian, Godfrey. 2023. "ChatGPT costs \$700,000 to run daily, OpenAI may go bankrupt in 2024- report." *TechNext*. August 14. <https://technext24.com/2023/08/14/chatgpt-costs-700000-daily-openai/> (검색일: 2024. 1. 16.).

- Fetter, Steve. 1991. "Ballistic Missiles and Weapons of Mass Destruction: What Is the Threat? What Should be Done?" *International Security* 16(1): 5-41.
- Freedman, Lawrence. 2003. *The Evolution of Nuclear Strategy*. London: Palgrave Macmillan London.
- Gilpin, Robert. 1981. *War and Change in World Politics*. Cambridge: Cambridge University Press.
- Government of the Netherlands. 2023. "REAIM Call to Action." Summit on Responsible Artificial Intelligence in the Military Domain. February 16.
- Haenlein, Michael and Andreas Kaplan. 2019. "A Brief History of Artificial Intelligence: On the Past, Present, and Future of Artificial Intelligence." *California Management Review* 61(4): 5-14.
- Horowitz, Michael C. 2018. "Artificial Intelligence, International Competition, and the Balance of Power." *Texas National Security Review* 1(3). May.
- Jervis, Robert. 1978. "Cooperation Under the Security Dilemma." *World Politics* 30(2): 167-214.
- _____. 1989. *The Meaning of the Nuclear Revolution*. New York: Cornell University Press.
- Johnson, James. 2019. "Artificial intelligence & future warfare: implications for international security." *Defense & Security Analysis* 35(2): 147-169.
<https://doi.org/10.1080/14751798.2019.1600800>.
- Johnson, James. 2023. *AI and the Bomb: Nuclear Strategy and Risk in the Digital Age*. Oxford University Press.
- Kaplan, Fred. 1983. *The Wizards of Armageddon*. New York: Simon and Schuster.
- Kahn, Herman. 1960. *On Thermonuclear War*. Princeton University Press.
- Kissinger, Henry A. and Graham Allison. 2023. "The Path to AI Arms Control." *Foreign Affairs*. October 13. <https://www.foreignaffairs.com/united-states/henry-kissinger-path-artificial-intelligence-arms-control>
- Lindsay, Jon R. 2023/2024. "War Is from Mars, AI Is from Venus: Rediscovering the Institutional Context of Military Automation." *Texas National Security Review* 7(1): 29-47.

- Lushenko, Paul. 2023. "AI and the future of warfare: The troubling evidence from the US military." *Bulletin of the Atomic Scientists*. November 29.
- Metz, Steven. 2000. *Armed Conflict in the 21st Century: The Information Revolution and Post-Modern Warfare*. Carlisle Barracks: Strategic Studies Institute.
- Modelski, G., Thompson, W.R. 1988. "The Long Cycle of World Leadership." In *Seapower in Global Politics, 1494-1993*. London: Palgrave Macmillan.
https://doi.org/10.1007/978-1-349-09154-6_5.
- Pape, Robert A. 1996. *Bombing to Win: Air Power and Coercion in War*. Ithaca: Cornell University Press.
- Schneider, Jacquelyn and Julia Macdonald. 2024. "Looking back to look forward: Autonomous systems, military revolutions, and the importance of cost." *Journal of Strategic Studies* 47(2): 162-184. <https://doi.org/10.1080/01402390.2022.2164570>
- Snyder, Jack. 1984. "Civil-Military Relations and the Cult of the Offensive, 1914 and 1984." *International Security* 9(1): 108-146.
- Talmadge, Caitlin. 2017. "Would China Go Nuclear? Assessing the Risk of Chinese Nuclear Escalation in a Conventional War with the United States." *International Security* 41(4): 50-92.
- U.S. Department of Defense. 2018. "Summary of the 2018 Department of Defense Artificial Intelligence Strategy." <https://media.defense.gov/2019/Feb/12/2002088963/-1/-1/1/SUMMARY-OF-DOD-AISTRATEGY.PDF>. (검색일: 2024. 1. 16.).
- _____. 2023. "Data, Analytics, and Artificial Intelligence Adoption Strategy: Accelerating Decision Advantage." July 27. https://media.defense.gov/2023/Nov/02/2003333300/-1/-1/1/DOD_DATA_ANALYTICS_AI_ADOPTION_STRATEGY.PDF. (검색일: 2024. 1. 16.)
- Wagner, R. Harrison. 1991. "Nuclear Deterrence, Counterforce Strategies, and the Incentive to Strike First." *American Political Science Review* 85(3): 727-750.
- Waltz, Kenneth N. 2003. "Waltz Responds to Sagan." In Scott D. Sagan and Kenneth N. Waltz, *The Spread of Nuclear Weapons: A Debate Renewed*. New York: W.W. Norton & Company.

Wolfson, Richard and Ferenc Dalnoki-Veress. 2022. "The Devastating Effects of Nuclear Weapons." The MIT Press Reader. March 2.

김양규. 2023. "‘전장으로서의 우주’와 미래의 억지전략: 미중 안보전략 변화와 한국." 『미래전의 도전과 항공우주산업』. 신범식 편. 서울: 사회평론아카데미.

김양규. 2024. "2024 년 세계 군사질서와 한국: 정확성과 투명성 혁명에 따른 공격 우위 시대 한국의 안보정책." 신년기획 특별논평 시리즈 3. 11 월 20 일.

<https://www.eai.or.kr/new/ko/pub/view.asp?intSeq=22234> (검색일: 2024. 6. 2.).

■ **저자:** 김양규_EAI 수석연구원. 서울대학교 정치외교학부 강사.

■ **담당 및 편집:** 박지수_EAI 연구원

문의: 02-2277-1683 (ext. 208) jspark@eai.or.kr

본 대담록을 인용할 때에는 반드시 출처를 밝혀주시기 바랍니다.

EAI는 어떠한 정파적 이해와도 무관한 독립 연구기관입니다.

EAI가 발행하는 보고서와 저널 및 단행본에 실린 주장과 의견은 EAI와는 무관하며 오로지 저자 개인의 견해를 밝힙니다.

발행일 2024년 9월 6일

“[AI와 신문명 표준 스페셜리포트] 군사도전①: 인공지능-핵무기 넥서스(AI-Nuclear Nexus)와 세계군사질서 전망”
979-11-6617-796-5 95340

재단법인 동아시아연구원

03028 서울특별시 종로구 사직로7길 1
Tel. 82 2 2277 1683 Fax 82 2 2277 1684

Email eai@eai.or.kr

Website www.eai.or.kr