



ARTIFICIAL INTELLIGENCE NEGATIVE - DUDL -

ADAPTED FROM NAUDL CORE FILES



ARTIFICIAL INTELLIGENCE NEG

Artificial Intelligence Neg	3
Argument Glossary	5
Overviews	7
Overview (Affirmative)	7
Strategic Overview – China/Russia	8
Strategic Overview - Regulations	9
1NC Shell	10
Inherency	10
Regulations	11
China	15
Russia	19
LAWs used by Russia are lagging – need significant data training	20
Solvency	23
2NC	26
AT: “LAWs are good for military operations”	27
Regulations Ext.	33
US has a history of misusing AI and spurring mistrust in the technology, Flint proves	33
Solvency Ext.	42
China Ext.	50
Russia/Ukraine Ext.	52
More Neg Answers	58
Impact - Nuclear War	58



ARTIFICIAL INTELLIGENCE NEG

The goal of the negative is simple: Negate the plan of action presented by the affirmative, explain why the affirmative is a bad idea, and present new arguments which serve as reasons why the affirmative will make the status quo WORSE. If the judge believes the plan will not work or make the status quo worse, then neg wins

Speech Time	(Minutes)
1 st Affirmative Constructive (1AC)	8
2 nd Negative Speaker Questions 1 st Affirmative Speaker	3
1 st Negative Constructive (1NC)	8
1 st Affirmative Speaker Questions 1 st Negative Speaker	3
2 nd Affirmative Constructive (2AC)	8
1 st Negative Speaker Questions 2 nd Affirmative Speaker	3
2 nd Negative Constructive (2NC)	8
2 nd Affirmative Speaker Questions 2 nd Negative Speaker	3
1 st Negative Rebuttal (1NR)	5
1 st Affirmative Rebuttal (1AR)	5
2 nd Negative Rebuttal (Closing Statement) (2NR)	5
2 nd Affirmative Rebuttal (Closing Statement) (2AR)	5

Speaking Roles on the Negative:

- **1st Negative Speaker:** Your job is to introduce a range of negative arguments in the 1NC, and to definitively win at least one of those arguments in the 1NR.
- **2nd Negative Speaker:** Your job is to expand upon one or two arguments made in the 1NC, then to choose the best argument made by the negative team and show why the negative should win the debate in the 2NR. You are in charge of choosing negative strategy, since you'll have to explain it in the 2NR

Phases of a Debate:

1. **1NC:** Outline a few different reasons why the affirmative is a bad idea, without going into too much detail on any one of them.
2. **2NC/1NR:** Think of these as a single speech, given by different people. Each debater should choose one or two (different) arguments from the 1NC and go into greater detail, explaining and adding evidence when needed.
3. **2NR:** The second negative speaker should give a closing argument all about the strongest negative position (after hearing the affirmative speak in the 1AR). Tell the judge why the negative team should win.

Each negative case will have four main parts. You'll have to win each piece in order to win a debate as the negative.

- **Solvency:** How will the plan work? Will the affirmative solve what it is they are attempting to solve with their advantages? Will there be a different consequence if the affirmative is passed? In the context of the AI negative file, will cooperation with NATO truly solve Russian/China aggression? Will this trigger conflict with other nations?
- **Case impacts/defense:** These are arguments saying the impacts the affirmative claims to solve are not correct.
- **Case turns:** These arguments are negative consequences if the affirmative were to pass. In this instance, will US cooperation with NATO, in turn, trigger an arms race with China and Russia, inciting international conflict and violence?
- **Impact Framing:** These arguments answer the moral obligation claims the affirmative makes. How long will it take for the affirmative impacts to be solved? How many people does the affirmative claim to aid? How probable are the affirmative impacts?

National Association for Urban Debate Leagues



You and your partner will have a chance to support your arguments and read more evidence in the block (13-minute stretch of neg speeches, i.e. 2NC + 1NR).

Topic Introduction:

The artificial intelligence (AI) affirmative offers that we increase cooperation with the North Atlantic Treaty Organization (NATO) to improve/develop AI used in military operations. The main focus of this affirmative's AI technology is development of Lethal Autonomous Weapons (LAWs). Currently, the US has a limited cooperation with NATO in regards to LAWs. The affirmative introduces this cooperation with benefits that include updating outdated regulations on LAWs, a stronghold alliance against historically aggressive nations such as Russia/China, and the ability to curtail armed conflict and war through Russian escalation.

The Biden Administration and NATO have acknowledged the violence AI-powered autonomous weapons play in international conflicts without regulation or oversight. The affirmative argues they may effectively solve these issues via the plan, and cooperation with NATO with that concentration is the mechanism to do so.



NATO (noun) – also known as the North Atlantic Treaty Organization, an international alliance of 30 countries, which includes the United States.

Artificial Intelligence (noun) – computer programs that utilize machine learning and are able to mimic human intelligence

Cybersecurity (noun) – measures taken to protect a computer or computer system (as on the Internet) against unauthorized access or attack

Cyberattack (noun) – an attempt to gain illegal access to a computer or computer system for the purpose of causing damage or harm

Hegemony (noun) – a term in international relations used to describe a country has more power in economics, and military strength relative to others.

Counter intelligence (noun) – organized activity of an intelligence service designed to block an enemy's sources of information, to deceive the enemy, to prevent sabotage, and to gather political and military information

Autonomous Weapons (noun) – weapons which have the capability of functioning at some level without human input or supervision

Technological singularity (noun) – The idea of technological singularity, and what it would mean if ordinary human intelligence were enhanced or overtaken by artificial intelligence.

Atomation (noun) – Jobs being performed by Artificial intelligence

Superintelligence (noun) – AI with the capability to evolve on its own

Dexterous (adjective) – done with mental or physical skill, quickness, or grace; done with dexterity

Malicious (adjective) – An adjective that describes bad intent

Socioeconomic inequality (noun) – socioeconomic inequality is not measured by an arbitrary income line below which the poor are placed, but rather by the distances between the relative positions occupied by the various segments of society

Algorithmic trading (noun) – Algorithmic trading is a process for executing orders utilizing automated and pre-programmed trading instructions to account for variables such as price, timing and volume.

Machine learning (noun) – The capability for AI to consistently improve due to past trial and error

Authoritarian regime (noun) – A type of government that maintains near absolute control typically by force, showing little concern for public opinion, and governed by a single individual, group, or class.

Algorithmic bias (noun) – Humans who are biased develop AI, and because of this human bias is transferable to algorithms



National Association for Urban Debate Leagues

Strategic Concept 2022 (noun) – The Strategic Concept is a key document for the Alliance. It reaffirms NATO's values and purpose, and provides a collective assessment of the security environment. It also drives NATO's strategic adaptation and guides its future political and military development.

DIANA initiative (noun) – A diversity-driven conference committed to helping all underrepresented people in Information Security. The Diana Initiative features multiple speaker tracks, villages with hands-on workshops, and a Capture the Flag event.

Political capital (noun) – Pursuing tech leadership in order to advance political capital

Proliferation (noun) – Destruction of states by nuclear means



OVERVIEW (AFFIRMATIVE)

THE AFFIRMATIVE VERSION OF THIS FILE OUTLINES ACTION BY THE UNITED STATES FEDERAL GOVERNMENT (USFG) TO INCREASE COOPERATION WITH NATO ON DEVELOPMENT AND REGULATION OF ARTIFICIAL INTELLIGENCE, WITH CONCENTRATIONS IN LETHAL AUTONOMOUS WEAPONS (LAWs) USED IN MILITARY OPERATIONS.

While the affirmative may choose to run 1, 2, or all advantages in the 1AC, it is important to strategize and prepare the negative in a way that best responds to either scenario.

The affirmative can be run in a number of ways outlined below:

1. “Hard right” – big impacts like nuclear war, extinction, global catastrophe, etc.

- a. Russia advantage
- b. China advantage

This version of the affirmative will focus on the pursuit of technological leadership between Us, Russia, and China, and how Russia/China will find shortcuts, fuel the arms race towards AI advancement, and spur conflict and global war if they advance sooner/better than the US

2. “Soft left” – impacts which affect people’s quality of life and experience in a more realistic/fathomable way such was racial/gender discrimination, ageism, ableism, targeted policing of marginalized groups, etc.

- a. Regulations advantage

This version of the affirmative will focus on the regulations and ethics of LAWs and the use of artificial intelligence in this technology. The affirmative will outline the potential for misuse, discriminatory bias, and overdevelopment of AI leading to dangerous “killer robot” scenarios



THE NEGATIVE VERSION OF THIS FILE HAS A FEW DIFFERENT AVENUES TO TAKE, PENDING THE WAY THE AFFIRMATIVE IS CONSTRUCTED IN THE ROUND. A LOT OF THE EVIDENCE IS CENTERED AROUND THE EFFICACY IN AI AND WHY ITS DEVELOPMENT WILL SPUR INTERNATIONAL CONFLICT. NEGATIVE IMPACTS INCLUDE PROLIFERATION OF AI/LAWS, GREAT POWER WARS, CONFLICTS WITH RUSSIA/CHINA IN TECH LEADERSHIP COMPETITION, AND BIAS IN THE EXECUTION OF AI-POWERED TECHNOLOGY. BREAKDOWN BELOW:

1. CHINA ADVANTAGE

- a. China is currently lagging behind in AI research and development (R&D), but is currently pursuing tech leadership by 2030. The Creemers evidence explains this and how China will find shortcuts and early use of AI if it means the opportunity to surpass the US and other NATO countries. This will directly clash with the affirmative's Franke (21) evidence which expects US cooperation with NATO to result in the US being the leader in AI tech
- b. As in the past, China's focus will be advancing AI in the private sector. R&D of this technology will be dominated by companies looking to make a profit and will shortcut the advancement of AI which will prove detrimental if used in military operations
- c. Tensions already growing between US and China over investments in this kind of technology. Both nations will expand their advancement to include improved LAWS to increase "national security" and spur an arms race which can be catastrophic if faulty tech is used

2. RUSSIA ADVANTAGE

- a. Much like China, Russia lags in R&D of AI. However, as cooperation increases, Russia will pursue leadership in the industry to bolster their military. This proves incredibly problematic given the conflict between Russia and Ukraine. Russia has already used semi-autonomous drones in Ukraine and it is only a matter of time until they up their game to include LAWS – that's the Allen 5/26 evidence
 - i. This argument will directly clash with the affirmative's Laird (20) evidence which states cooperation to advance AI will foster diplomacy and reduce Russian aggression
- b. Putin's goal is obtaining the "military edge" in AI tech. The proliferation of this LAWS ensures Putin will use his arsenal against nations where conflict arises, even if shortcuts were made to advance this technology



THE NEGATIVE VERSION OF THIS FILE HAS A FEW DIFFERENT AVENUES TO TAKE, PENDING THE WAY THE AFFIRMATIVE IS CONSTRUCTED IN THE ROUND. A LOT OF THE EVIDENCE IS CENTERED AROUND THE EFFICACY IN AI AND WHY ITS DEVELOPMENT WILL SPUR INTERNATIONAL CONFLICT. NEGATIVE IMPACTS INCLUDE PROLIFERATION OF AI/LAWS, GREAT POWER WARS, CONFLICTS WITH RUSSIA/CHINA IN TECH LEADERSHIP COMPETITION, AND BIAS IN THE EXECUTION OF AI-POWERED TECHNOLOGY. BREAKDOWN BELOW:

1. REGULATIONS ADVANTAGE

- a. The affirmative is focused on AI-powered lethal autonomous weapons. However, some of the evidence in this file will take a broader approach in what is known as “generics.” These generics outline the broader issue of artificial intelligence itself and the issues with its programming in the context of biases, inaccuracies, and unreliability
 - i. The generics mentioned above may also be connected to LAWS, their harmful consequences, and the efficacy of this technology in some examples below
- b. Some of the generics describe how AI technology has historically targeted marginalized communities via automation of the workforce.
 - i. There is an important application in standardizing this statement to include automation of military workforce since there will be less personnel to operate drones as advancement will ensure weapons are fully autonomous
- c. There is also a portion of evidence which can be applied from the regulations advantage to both Russia/China advantages. This evidence argues that as an arms race unfolds and US/Russia/China pursue this leadership, the focus will shift from military applications to making profit off the private sector in selling this tech – Thomas (21) evidence
 - i. This will directly clash with the affirmative’s Stanley-Lockman (21) evidence because the affirmative claims “values are at the foundation of US alliances.” The argument the negative can expand upon is saying the affirmative will be so focused on profiting off of AI technology, it will begin to overlook shortcuts and the inherent issues presented in the previous points
- d. The last set of arguments brings to light the bias in creation and development of AI
 - i. Statistically speaking, most R&D of AI is conducted by older white men – Thomas (21) evidence to quote
 - ii. This means there is a lack of diversity in the development and creation of this technology which can prove un/intentionally discriminatory in its application
 - iii. AI regulation will fail unless there is a SPECTRUM of identities in its developing and programming to include marginalized groups from all over the world
 1. This argument may also be applied to the solvency to dispute the efficacy of the USFG being the best actor to cooperate with NATO



INHERENCY

Responds to Michaelson 21 – regulations not needed

AI DATA HANDLING WEAK NOW AND PLAN IS NON-UNIQUE — EFFORTS LIKE THE DIANA INITIATIVE ALREADY IN MOTION TO IMPROVE AI R&D

Dolan, 6/8 -- Dolan, C. (2022, June 8). NATO's 2022 strategic concept must enhance digital access and capacities. Just Security. Retrieved June 12, 2022, from <https://www.justsecurity.org/81839/natos-2022-strategic-concept-must-enhance-digital-access-and-capacities/>

Essential Role of Artificial Intelligence The challenge for NATO is not necessarily adopting and investing in emerging and disrupting technologies for collective defense. Rather, the question is whether ACT and ACO can enhance accessibility to digital platforms and ease communications between platforms. Here, artificial intelligence (AI) can play a role in overcoming critical obstacles. AI is now occupying a greater space in NATO's collective defense orientation. The challenge in the 2022 Strategic Concept will be delineating the degree to which AI will enhance the ability of the alliance to analyze information and assess data. Moreover, AI is only as good as the data it relies on. To maintain its technological edge, in 2021 NATO released an Artificial Intelligence Strategy, a good step toward maximizing interoperability of weapons systems, improving infrastructure, and building resilient hybrid defenses. While it emphasizes collaboration with the private sector and academia, the strategy needs further refinement as AI would help NATO's military and civilian personnel interlink devices on different platforms, perform rigorous data analytics, and quicken response time in response to a conventional or hybrid attack. One innovation is the Defence Innovation Accelerator for the North Atlantic (DIANA). DIANA leverages partnerships among academics, technology companies, start-up firms to address the full spectrum of threats in the security environment. Private sector firms utilizing DIANA can access innovation sites and test centers that focus on artificial intelligence, machine learning, quantum computing, autonomous machines, and biotechnology. In the past, alliance members supported the creation and development of a venture capital Innovation Fund for technology companies and start-ups. The Strategic Concept should support the initiative with sustained public funds in ways that allow for strategic-level planning at ACO and ATO to communicate more effectively at the operational and tactical levels. DIANA and the Innovation Fund are models for the alliance. An existing partnership that is proving effective is NATO's collaboration with Klarrio, a firm that provides the alliance with innovative real-time data analytics and streaming services to combat disinformation and fake news. Klarrio has partnered with NATO's Strategic Communications Center (StratCom) to assist the alliance in the information domain by streaming data analysis, processing, and applications analytics to identify and eradicate disinformation for NATO's strategic-level planners. It supports and updates dashboard services to track suspicious activities in social media platforms, maintains an interactive user interface, generates analytical reports, and uses machine learning to analyze data and information.



AI POSES LAUNDRY LIST OF ISSUES — ADVANCEMENT BEYOND CONTROL CANNOT BE DISMISSED

Thomas, 21 -- Thomas, M. (2021, July 7). *7 dangerous risks of Artificial Intelligence*. Built In. Retrieved June 8, 2022, from <https://builtin.com/artificial-intelligence/risks-of-artificial-intelligence>

In March 2018, at the South by Southwest tech conference in Austin, Texas, Tesla and SpaceX founder Elon Musk issued a friendly warning: “Mark my words,” he said, billionaire casual in a furry-collared bomber jacket and days’ old scruff, “AI is far more dangerous than nukes.” No shrinking violet, especially when it comes to opining about technology, the outspoken Musk has repeated a version of these artificial intelligence premonitions in other settings as well. “I am really quite close... to the cutting edge in AI, and it scares the hell out of me,” he told his SXSW audience. “It’s capable of vastly more than almost anyone knows, and the rate of improvement is exponential.” RISKS OF ARTIFICIAL INTELLIGENCE Automation-spurred job loss Privacy violations ‘Deepfakes’ Algorithmic bias caused by bad data Socioeconomic inequality Market volatility Weapons automatization Musk, though, is far from alone in his exceedingly skeptical (some might say bleakly alarmist) views. A year prior, the late physicist Stephen Hawking was similarly forthright when he told an audience in Portugal that AI’s impact could be cataclysmic unless its rapid development is strictly and ethically controlled. “Unless we learn how to prepare for, and avoid, the potential risks,” he explained, “AI could be the worst event in the history of our civilization.” Considering the number and scope of unfathomably horrible events in world history, that’s really saying something. And in case we haven’t driven home the point quite firmly enough, research fellow Stuart Armstrong from the Future of Life Institute has spoken of AI as an “extinction risk” were it to go rogue. Even nuclear war, he said, is on a different level destruction-wise because it would “kill only a relatively small proportion of the planet.” Ditto pandemics, “even at their more virulent.” Musk discusses his fear of AI **“If AI went bad, and 95 percent of humans were killed,” he said, “then the remaining five percent would be extinguished soon after.** So despite its uncertainty, it has certain features of very bad risks.” How, exactly, would AI arrive at such a perilous point? Cognitive scientist and author Gary Marcus offered some details in an illuminating 2013 New Yorker essay. The smarter machines become, he wrote, the more their goals could shift. **“Once computers can effectively reprogram themselves, and successively improve themselves, leading to a so-called ‘technological singularity’ or ‘intelligence explosion,’ the risks of machines outwitting humans in battles for resources and self-preservation cannot simply be dismissed.”** Is Artificial Intelligence a Threat? As AI grows more sophisticated and ubiquitous, the voices warning against its current and future pitfalls grow louder. Whether it’s the increasing automation of certain jobs, gender and racial bias issues stemming from outdated information sources or autonomous weapons that operate without human oversight (to name just a few), unease abounds on a number of fronts. And we’re still in the very early stages. IS ARTIFICIAL INTELLIGENCE A THREAT? The tech community has long-debated the threats posed by artificial intelligence. **Automation of jobs, the spread of fake news and a dangerous arms race of AI-powered weaponry have been proposed as a few of the biggest dangers posed by AI.** Destructive superintelligence — aka artificial general intelligence that’s created by humans and escapes our control to wreak havoc — is in a category of its own. It’s also something that might or might not come to fruition (theories vary), so at this point it’s less risk than hypothetical threat — and ever-looming source of existential dread. Here are some of the ways artificial intelligence poses a serious risk: JOB AUTOMATION Job automation is generally viewed as the most immediate concern. **It’s no longer a matter of if AI will replace certain types of jobs, but to what degree.**



AUTOMATION OF THE WORKFORCE WITH THE USE OF AI TARGETS MARGINALIZED GROUPS. JOB LOSS WILL BE INEVITABLE WITH THIS ADVANCEMENT

Thomas, 21 -- Thomas, M. (2021, July 7). *7 dangerous risks of Artificial Intelligence*. Built In. Retrieved June 8, 2022, from <https://builtin.com/artificial-intelligence/risks-of-artificial-intelligence>

In many industries — particularly but not exclusively those whose workers perform predictable and repetitive tasks — disruption is well underway. According to a 2019 Brookings Institution study, 36 million people work in jobs with “high exposure” to automation, meaning that before long at least 70 percent of their tasks — ranging from retail sales and market analysis to hospitality and warehouse labor — will be done using AI. An even newer Brookings report concludes that white collar jobs may actually be most at risk. And per a 2018 report from McKinsey & Company, the African American workforce will be hardest hit. “The reason we have a low unemployment rate, which doesn’t actually capture people that aren’t looking for work, is largely that lower-wage service sector jobs have been pretty robustly created by this economy,” renowned futurist Martin Ford told Built In. “I don’t think that’s going to continue.” As AI robots become smarter and more dexterous, he added, the same tasks will require fewer humans. And while it’s true that AI will create jobs, an unspecified number of which remain undefined, many will be inaccessible to less educationally advanced members of the displaced workforce. “If you’re flipping burgers at McDonald’s and more automation comes in, is one of these new jobs going to be a good match for you?” Ford said. “Or is it likely that the new job requires lots of education or training or maybe even intrinsic talents — really strong interpersonal skills or creativity — that you might not have? Because those are the things that, at least so far, computers are not very good at.” John C. Havens, author of *Heartificial Intelligence: Embracing Humanity and Maximizing Machines*, calls bull on the theory that AI will create as many or more jobs than it replaces. About four years ago, Havens said, he interviewed the head of a law firm about machine learning. The man wanted to hire more people, but he was also obliged to achieve a certain level of returns for his shareholders. A \$200,000 piece of software, he discovered, could take the place of ten people drawing salaries of \$100,000 each. That meant he’d save \$800,000. The software would also increase productivity by 70 percent and eradicate roughly 95 percent of errors. From a purely shareholder-centric, single bottom-line perspective, Havens said, “there is no legal reason that he shouldn’t fire all the humans.” Would he feel bad about it? Of course. But that’s beside the point. Even professions that require graduate degrees and additional post-college training aren’t immune to AI displacement. In fact, technology strategist Chris Messina said, some of them may well be decimated. AI already is having a significant impact on medicine. Law and accounting are next, Messina said, the former being poised for “a massive shakeup.” “Think about the complexity of contracts, and really diving in and understanding what it takes to create a perfect deal structure,” he said. “It’s a lot of attorneys reading through a lot of information — hundreds or thousands of pages of data and documents. It’s really easy to miss things. So AI that has the ability to comb through and comprehensively deliver the best possible contract for the outcome you’re trying to achieve is probably going to replace a lot of corporate attorneys.” Accountants should also prepare for a big shift, Messina warned. Once AI is able to quickly comb through reams of data to make automatic decisions based on computational interpretations, human auditors may well be unnecessary. **PRIVACY, SECURITY AND THE RISE OF ‘DEEPFAKES’** While job loss is currently the most pressing issue related to AI disruption, it’s merely one among many potential risks. In a February 2018 paper titled “The Malicious Use of Artificial Intelligence:



Thomas, 21 -- Thomas, M. (2021, July 7). 7 dangerous risks of Artificial Intelligence. Built In. Retrieved June 8, 2022, from <https://builtin.com/artificial-intelligence/risks-of-artificial-intelligence>

Forecasting, Prevention, and Mitigation,” 26 researchers from 14 institutions (academic, civil and industry) enumerated a host of other dangers that could cause serious harm — or, at minimum, sow minor chaos — in less than five years. **“Malicious use of AI,”** they wrote in their 100-page report, **“could threaten digital security** (e.g. through criminals training machines to hack or socially engineer victims at human or superhuman levels of performance), **physical security** (e.g. non-state actors weaponizing consumer drones), and **political security** (e.g. through privacy-eliminating surveillance, profiling, and repression, or through automated and targeted disinformation campaigns).” In addition to its more existential threat, Ford is focused on the way AI will adversely affect privacy and security. A prime example, he said, is China’s “Orwellian” use of facial recognition technology in offices, schools and other venues. But that’s just one country. “A whole ecosphere” of companies specialize in similar tech and sell it around the world. What we can so far only guess at is whether that tech will ever become normalized. As with the internet, where we blithely sacrifice our digital data at the altar of convenience, will round-the-clock, AI-analyzed monitoring someday seem like a fair trade-off for increased safety and security despite its nefarious exploitation by bad actors? “Authoritarian regimes use or are going to use it,” Ford said. “The question is, How much does it invade Western countries, democracies, and what constraints do we put on it?” “Authoritarian regimes use or are going to use it ... The question is, How much does it invade Western countries, democracies, and what constraints do we put on it?” **AI will also give rise to hyper-real-seeming social media “personalities” that are very difficult to differentiate from real ones,** Ford said. **Deployed cheaply and at scale on Twitter, Facebook or Instagram, they could conceivably influence an election.** The same goes for so-called audio and video deepfakes, created by manipulating voices and likenesses. The latter is already making waves. But the former, Ford thinks, will prove immensely troublesome. **Using machine learning, a subset of AI that’s involved in natural language processing, an audio clip of any given politician could be manipulated to make it seem as if that person spouted racist or sexist views** when in fact they uttered nothing of the sort. If the clip’s quality is high enough so as to fool the general public and avoid detection, Ford added, **it could “completely derail a political campaign.”** And all it takes is one success. From that point on, he noted, “no one knows what’s real and what’s not. So it really leads to a situation where you literally cannot believe your own eyes and ears; you can’t rely on what, historically, we’ve considered to be the best possible evidence... That’s going to be a huge issue.” Lawmakers, though frequently less than tech-savvy, are acutely aware and pressing for solutions. **AI BIAS AND WIDENING SOCIOECONOMIC INEQUALITY** **Widening socioeconomic inequality sparked by AI-driven job loss is another cause for concern. Along with education, work has long been a driver of social mobility.** However, when it’s a certain kind of work — the predictable, repetitive kind that’s prone to AI takeover — research has shown that those who find themselves out in the cold are much less apt to get or seek retraining compared to those in higher-level positions who have more money. (Then again, not everyone believes that.) **Various forms of AI bias are detrimental, too.** Speaking recently to the New York Times, Princeton computer science professor Olga Russakovsky said it goes well beyond gender and race. In addition to data and algorithmic bias (the latter of which can “amplify” the former), AI is developed by humans and humans are inherently biased. **“A.I. researchers are primarily people who are male, who come from certain racial demographics, who grew up in high socioeconomic areas, primarily people without disabilities,”**



AI REGULATION **MUST** INCLUDE AN ARRAY OF EXPERTS ACROSS A SPECTRUM OF IDENTITIES

Thomas, 21 -- Thomas, M. (2021, July 7). *7 dangerous risks of Artificial Intelligence*. Built In. Retrieved June 8, 2022, from <https://builtin.com/artificial-intelligence/risks-of-artificial-intelligence>

Mitigating the Risks of AI Many believe the only way to prevent or at least temper the most malicious AI from wreaking havoc is some sort of regulation. “I am not normally an advocate of regulation and oversight — I think one should generally err on the side of minimizing those things — but this is a case where you have a very serious danger to the public,” Musk said at SXSW. “It needs to be a public body that has insight and then oversight to confirm that everyone is developing AI safely. This is extremely important.” Ford agrees — with a caveat. Regulation of AI implementation is fine, he said, but not of the research itself. “You regulate the way AI is used,” he said, “but you don’t hold back progress in basic technology. I think that would be wrong-headed and potentially dangerous.” “You regulate the way AI is used ... but you don’t hold back progress in basic technology. I think that would be wrong-headed and potentially dangerous.” Because any country that lags in AI development is at a distinct disadvantage — militarily, socially and economically. The solution, Ford continued, is selective application: “We decide where we want AI and where we don’t; where it’s acceptable and where it’s not. And different countries are going to make different choices. So China might have it everywhere, but that doesn’t mean we can afford to fall behind them in the state-of-the-art.” Speaking about autonomous weapons at Princeton University, American General John R. Allen emphasized the need for “a robust international conversation that can embrace what this technology is.” If necessary, he went on, there should also be a conversation about how best to control it, be that a treaty that fully bans AI weapons or one that permits only certain applications of the technology. For Havens, safer AI starts and ends with humans. His chief focus, upon which he expounds in his 2016 book, is this: “How will machines know what we value if we don’t know ourselves?” In creating AI tools, he said, it’s vitally important to “honor end-user values with a human-centric focus” rather than fixating on short-term gains. “Technology has been capable of helping us with tasks since humanity began,” Havens wrote in *Heartificial Intelligence*. “But as a race we’ve never faced the strong possibility that machines may become smarter than we are or be imbued with consciousness. This technological pinnacle is an important distinction to recognize, both to elevate the quest to honor humanity and to best define how AI can evolve it. That’s why we need to be aware of which tasks we want to train machines to do in an informed manner. This involved individual as well as societal choice.” The World Economic Forum discusses implementing responsible AI AI researchers Fei-Fei Li and John Etchemendy, of Stanford University’s Institute for Human-Centered Artificial Intelligence, feel likewise. In a recent blog post, they proposed involving numerous people in an array of fields to make sure AI fulfills its huge potential and strengthens society instead of weakening it: “Our future depends on the ability of social- and computer scientists to work side-by-side with people from multiple backgrounds — a significant shift from today’s computer science-centric model,” they wrote. “The creators of AI must seek the insights, experiences and concerns of people across ethnicities, genders, cultures and socio-economic groups, as well as those from other fields, such as economics, law, medicine, philosophy, history, sociology, communications, human-computer-interaction, psychology, and Science and Technology Studies (STS). This collaboration should run throughout an application’s lifecycle — from the earliest stages of inception through to market introduction and as its usage scales.”



PURSUIT OF CAPITAL FROM AI DEVELOPMENT MAKES ARMS RACE INEVITABLE

Thomas, 21 -- Thomas, M. (2021, July 7). *7 dangerous risks of Artificial Intelligence*. Built In. Retrieved June 8, 2022, from <https://builtin.com/artificial-intelligence/risks-of-artificial-intelligence>

Russakovsky said. “We’re a fairly homogeneous population, so it’s a challenge to think broadly about world issues.” In the same article, Google researcher Timnit Gebru said the root of bias is social rather than technological, and called scientists like herself “some of the most dangerous people in the world, because we have this illusion of objectivity.” The scientific field, she noted, “has to be situated in trying to understand the social dynamics of the world, because most of the radical change happens at the social level.” And technologists aren’t alone in sounding the alarm about AI’s potential socio-economic pitfalls. Along with journalists and political figures, Pope Francis is also speaking up — and he’s not just whistling Sanctus. At a late-September Vatican meeting titled, “The Common Good in the Digital Age,” Francis warned that AI has the ability to “circulate tendentious opinions and false data that could poison public debates and even manipulate the opinions of millions of people, to the point of endangering the very institutions that guarantee peaceful civil coexistence.” “If mankind’s so-called technological progress were to become an enemy of the common good,” he added, “this would lead to an unfortunate regression to a form of barbarism dictated by the law of the strongest.” **A big part of the problem, Messina said, is the private sector’s pursuit of profit above all else.** Because “that’s what they’re supposed to do,” he said. “And so they’re not thinking of, ‘What’s the best thing here? What’s going to have the best possible outcome?’” **“The mentality is, ‘If we can do it, we should try it; let’s see what happens.’”** **“And if we can make money off it, we’ll do a whole bunch of it.”** But that’s not unique to technology. That’s been happening forever.” **AUTONOMOUS WEAPONS AND A POTENTIAL AI ARMS RACE** Not everyone agrees with Musk that AI is more dangerous than nukes, including Ford. But what if AI decides to launch nukes — or, say, biological weapons — sans human intervention? Or, what if an enemy manipulates data to return AI-guided missiles whence they came? Both are possibilities. And both would be disastrous. The more than 30,000 AI/robotics researchers and others who signed an open letter on the subject in 2015 certainly think so. “The key question for humanity today is whether to start a global AI arms race or to prevent it from starting,” they wrote. **“If any major military power pushes ahead with AI weapon development, a global arms race is virtually inevitable, and the endpoint of this technological trajectory is obvious: autonomous weapons will become the Kalashnikovs of tomorrow.** Unlike nuclear weapons, they require no costly or hard-to-obtain raw materials, so they will become ubiquitous and cheap for all significant military powers to mass-produce. It will only be a matter of time until they appear on the black market and in the hands of terrorists, dictators wishing to better control their populace, warlords wishing to perpetrate ethnic cleansing, etc. Autonomous weapons are ideal for tasks such as assassinations, destabilizing nations, subduing populations and selectively killing a particular ethnic group. We therefore believe that a military AI arms race would not be beneficial for humanity. There are many ways in which AI can make battlefields safer for humans, especially civilians, without creating new tools for killing people.” (The U.S. Military’s proposed budget for 2020 is \$718 billion. Of that amount, nearly \$1 billion would support AI and machine learning for things like logistics, intelligence analysis and, yes, weaponry.) A story in Vox detailed a frightening scenario involving the development of a sophisticated AI system “with the goal of, say, estimating some number with high confidence. The AI realizes it can achieve more confidence in its calculation if it uses all the world’s computing hardware, and it realizes that releasing a biological superweapon to wipe out humanity would allow it free use of all the hardware.



CHINA

Responds to Franke 21 – strong NATO front means US becomes tech leader

CHINA WANTS TO BE AI TECH LEADER BY 2030 BUT LAGS IN PRODUCTIVITY

Creemers et al, 2017 -- Creemers, R., Triolo, P., Webster, G., & Kania, E. B. (2017, August 1). *China's plan to 'lead' in AI: Purpose, prospects, and Problems*. New America. Retrieved June 13, 2022, from <https://www.newamerica.org/cybersecurity-initiative/blog/chinas-plan-lead-ai-purpose-prospects-and-problems/>

China's Plan to 'Lead' in AI: Purpose, Prospects, and Problems The present global verve about artificial intelligence (AI) and machine learning technologies has resonated in China as much as anywhere on earth. With the State Council's issuance of the "New Generation Artificial Intelligence Development Plan" (AIDP, 新一代人工智能发展规划) on July 20, China's government set out an ambitious roadmap including targets through 2030. Meanwhile, in China's leading cities, flashy conferences on AI have become commonplace. It seems every mid-sized tech company wants to show off its self-driving car efforts, while numerous financial tech start-ups tout an AI-driven approach. Chatbot startups clog investors' date books, and Shanghai metro ads pitch AI-taught English language learning. View the full translation here. The surge of showmanship and investment in the private sector is paralleled by a new focus among officials and policy thinkers in academia, civilian and military sectors of government, and major companies. Questions about how to regulate AI, and how to develop and use it ethically, have become a major topic for China's digital policy brain trust in recent months. At this point, investors, policymakers, and even engineers might wonder where the hype ends and the reality of China's agenda for "AI 2.0" begins. This new development plan is certainly (and typically for an aspirational central government document) packed with vagaries and grandiose ambitions, as China declares its intentions to pursue a "first-mover advantage" to become "the world's primary AI innovation center" by 2030. Nonetheless, the plan contains real signals and measures worthy of attention. Its specific and nonspecific goals, its bureaucratic positioning, and its long time-horizon make it an important reference point for a wide variety of policy, business, and security developments in coming years. In this three-part brief, produced by a team of analysts with diverse backgrounds who have jointly translated the plan in full, we present three perspectives on the new document and China's AI development trajectory. In part one, Rogier Creemers of the Leiden Asia Centre puts the plan in the context of China's own framework for regulatory problem-solving. In part two, Paul Triolo of the Eurasia Group and Graham Webster of Yale Law School's Paul Tsai China Center describe the prospects for leadership in AI and China's policy world linked with the plan. And in part three, Elsa Kania addresses China's pursuit of indigenous innovation to enable its advances in next-generation AI. Part I: Not a Moonshot, but a Legacy of Central Planning By Rogier Creemers In order to understand the drivers for the drafting of the New Generation Artificial Intelligence Development Plan, it is instructive to position it within the logic of Chinese development planning. Chinese policymaking resembles a champagne pyramid, in which the topmost glass represents what is called the "dominant contradiction": the fundamental tension animating a particular historical phase. Analysis of this tension then produces a cascade of subordinate problems at various layers, all of which must be solved in order to defuse the dominant contradiction and move forward to the next stage. This intellectual construct is a persistent residue of the economic planning of yore. The dominant contradiction at present, according to China's central leadership, is the tension between the Chinese nation's high material needs and its comparatively lagging productivity. This manifests in both comparative material poverty, and weak and ineffectual government.



Creemers et al, 2017 -- Creemers, R., Triolo, P., Webster, G., & Kania, E. B. (2017, August 1). *China's plan to 'lead' in AI: Purpose, prospects, and Problems*. New America. Retrieved June 13, 2022, from <https://www.newamerica.org/cybersecurity-initiative/blog/chinas-plan-lead-ai-purpose-prospects-and-problems/>

Since the days of Deng Xiaoping, China has looked to science and technology as critical in overcoming these issues, leading to considerable efforts to acquire, borrow or sometimes outright purloin foreign know-how, as well as develop China's own innovative capability. In recent years, "indigenous innovation" has received greater emphasis in an effort to move up the global value chain. Reliance on foreign technology is also seen as a security risk (although China's case isn't helped by the fact that, for instance, it long condoned the widespread installation of unlicensed and therefore unsecured versions of Windows). In this process, China has progressed through various iterations of what it calls "informatization." In governmental affairs, this included, amongst others, making government data transparent, opening government websites and social media accounts, and expanding Internet access to enable the better provision of public services. This agenda has served multiple, sometimes conflicting objectives, including social and economic development, but also expanding the ability of the central government to obtain information about local goings-on. As a result, its implementation record is mixed. On the one hand, China's Internet penetration rate is growing rapidly, and some of its businesses have become phenomenally successful; on the other hand, a lack of focus on security and compatibility led to severe infrastructural and security issues that still linger today. The current AI agenda has one of the same problems as the mixed-objective informatization drive: It is not a moonshot. The objective of sending someone to the Moon and back safely had the advantage of simplicity and clarity. In comparison, the current plan looks rather like a Santa's list of desiderata and objectives, but with little insight into how these should be achieved other than by throwing money at the problem. Several particular sections, phrases, and terms are clearly included not for policy effect, but to soothe ruffled feathers within a complex bureaucracy where consensus is required. That being said, it would be a mistake to assume the creation of these plans is merely a paper-generating exercise. Real political capital is spent on drafting a document such as this, and considerable resources will be invested in its implementation. Moreover, it would also be an error to merely see this as a Chinese ploy for world domination. While analysts tend to look mostly at the way in which Chinese policies influence global affairs or the fates of "our" companies, the focus on the Chinese side lies squarely with addressing the enormous societal challenges that China knows it faces. While foreign observers have emphasized the role that AI may play in modernizing China's military, this is only one goal among others, for instance mitigating the exploding costs of healthcare in a rapidly aging country, providing more effective transportation in China's metropolitan areas, or making civil servants more accountable for their conduct. Part II: Visions of Leadership for China in AI, and for Tech in Chinese Politics By Paul Triolo and Graham Webster In laying out its top-line goals for economic development and AI, the new State Council plan declares that in under a decade, AI will become "the main driving force for China's industrial upgrading and economic transformation." Statements such as these underline the ambition captured in the plan itself, but as always in Chinese politics, attention is also due to whose interests and ambitions are driving a given agenda. Promising government investment, and recognizing where China lags. The plan prescribes a high level of government investment in theoretical and applied AI breakthroughs (see Part III below for more), while also acknowledging that, in China as around the world --



AI CAPABILITIES ARE VAST BUT CHINA'S FOCUS IS POLITICAL CAPITAL AND THE PRIVATE SECTOR

Creemers et al, 2017 -- Creemers, R., Triolo, P., Webster, G., & Kania, E. B. (2017, August 1). *China's plan to 'lead' in AI: Purpose, prospects, and Problems*. New America. Retrieved June 13, 2022, from <https://www.newamerica.org/cybersecurity-initiative/blog/chinas-plan-lead-ai-purpose-prospects-and-problems/>

-- private companies are currently leading the charge on commercial applications of AI. Large companies in cloud services, e-commerce, social media, or other sectors with access to large troves of data that can be used to train AI algorithms are naturally positioned to lead in a variety of fields, including facial recognition, voice recognition, and natural language processing. The plan acknowledges, meanwhile, that China remains far behind world leaders in development of key hardware enablers of AI, such as microchips suited for machine learning use (e.g., GPUs or re-configurable processors). The plan's ambition is underlined by its recognition of the hard road ahead. Reconciling advances in AI with risks of disruption. With the proliferation of AI, China's government recognizes that new risks and challenges will arise for governance, economic security, and social stability. The plan also focuses on minimizing these risks to ensure the "safe, reliable, and controllable" development of AI. The plan includes formulation of laws, regulations, and ethical norms on AI, as well as mechanisms for safety and supervision. The plan seeks to mitigate likely negative externalities, such as job losses, associated with AI, while fully leveraging the opportunities. The Political Layer Harnessing the rise of AI to keep science and technology bureaucrats relevant. The new plan is clearly driven by China's waning Ministry of Science and Technology (MOST) and related science and technology (S&T) bureaucracy, as it links progress in AI with virtually every other S&T program of note. In the plan, the a National S&T Structural Reform and Innovation Systems Construction Leading Small Group, one of the several coordinating bodies (including the sometimes competing "Cybersecurity and Informatization" group that sits above the Cyberspace Administration of China) through which President Xi Jinping exercises influence and power, is heralded as the driver of the AI program and particularly the AI-related legal and regulatory system that the plan calls for. In addition, the plan calls to create a new AI Plan Promotion Office within MOST. For China's government, another digital solution to provide public goods. One theme of the plan is that AI can serve as a vehicle through which the Chinese government can provide better governance to the Chinese people, using AI to drive smart cities, smart government, smart manufacturing, and forming the infrastructure for a smart society. According to the plan's lofty aspirations, AI applications in agriculture, transportation, social security, pension management, public security, and a host of other government functions will enable the government to provide new levels of service and benefit to the Chinese nation. Informatization-linked governance improvements are not a new idea for China. A decade ago, governance sections in big Chinese bookstores started to fill with future-oriented e-government titles, which are joined now by guides on modernizing government services through the earlier official "Internet Plus" concept, and now how to interact with the public through microblogs and Tencent's ubiquitous WeChat platform. Growing power for China's tech players? If the application of AI tools in government becomes as successful as this plan envisions, it could create a comprehensive new power source within China's Party-state architecture, similar to the way China's security apparatus did during the tenure of the now-imprisoned former chief Zhou Yongkang. Moreover, AI capabilities could increase the influence and wherewithal of the tech industry within Chinese society and politics. The AI plan is yet another signal validating the increasing flow of elite money into the high-tech sector generally and AI and automation-related ventures in particular.



CHINA'S RAPID AI ADVANCEMENT ALREADY CAUSING TENSIONS BETWEEN THEM AND US. THIS INCREASES SKEPTICISM OF CHINESE AI INVESTMENTS

Creemers et al, 2017 -- Creemers, R., Triolo, P., Webster, G., & Kania, E. B. (2017, August 1). *China's plan to 'lead' in AI: Purpose, prospects, and Problems*. New America. Retrieved June 13, 2022, from <https://www.newamerica.org/cybersecurity-initiative/blog/chinas-plan-lead-ai-purpose-prospects-and-problems/>

Consequently, the Chinese government is encouraging its own AI enterprises to pursue an approach of “going out,” including through overseas mergers and acquisitions, equity investments, and venture capital, as well as the establishment of research and development centers abroad. These measures will undoubtedly prove controversial in some quarters and could provoke further frictions. For instance, Chinese investments in Silicon Valley AI startups have fueled a current U.S. debate on whether to broaden the remit of the Committee on Foreign Investment in the United States (CFIUS) to expand reviews of Chinese high-tech investments, particularly in AI. However, in the longer term, China will likely become less dependent upon foreign innovation resources as its own capacity advances. Cognizant of these challenges, China plans to develop resources and ecosystems conducive to the goal of becoming a “premier innovation center” in AI science and technology by 2030. In support of this goal, the plan calls for an “open source and open” approach that takes advantage of synergies among industry, academia, research, and applications, including through creating AI “innovation clusters.” In practice, such an open and collaborative approach would involve the establishment of open-source, shared platforms and infrastructures for the enabling hardware and software. Certain tools and resources could be designed for common use by a variety of stakeholders and enterprises, including a data resource library, cloud service platforms, and methods to evaluate AI systems. Critically, the Chinese leadership and Chinese companies prioritize the attraction and development of leading talent in AI as a vital enabler of competitiveness. To achieve this, the Chinese government plans to continue to use a number of recruitment and talent programs, such as the “Thousand Talents” plan, which seeks to incentivize scientists in strategic domains to pursue their research in China. Concurrently, the Chinese government will focus on improving education and training in AI to strengthen its talent pool. In practice, this will entail the construction of a new AI discipline, including the establishment of AI majors in universities, the founding of dedicated AI institutes, and the expansion of enrollment in master’s and Ph.D. programs. The introduction of new higher education and vocational training programs will help to prepare China’s workforce for AI’s disruption of current employment structures.



RUSSIA

LAWs USED BY RUSSIA ARE LAGGING — NEED SIGNIFICANT DATA TRAINING

Allen, 5/26 -- Allen, G. C. (2022, May 26). *Russia probably has not used AI-enabled weapons in Ukraine, but that could change.* Russia Probably Has Not Used AI-Enabled Weapons in Ukraine, but That Could Change | Center for Strategic and International Studies. Retrieved June 13, 2022, from <https://www.csis.org/analysis/russia-probably-has-not-used-ai-enabled-weapons-ukraine-could-change>

In March, WIRED ran a story with the headline “Russia’s Killer Drone in Ukraine Raises Fears About AI in Warfare,” with the subtitle, “The maker of the lethal drone claims that it can identify targets using artificial intelligence.” The story focused on the KUB-BLA, a small kamikaze drone aircraft that smashes itself into enemy targets and detonates an onboard explosive. The KUB-BLA is made by ZALA Aero, a subsidiary of the Russian weapons manufacturer Kalashnikov (best known as the maker of the AK-47), which itself is partly owned by Rostec, a part of Russia’s government-owned defense-industrial complex. The WIRED story understandably attracted a lot of attention, but those who only read the sensational headline missed the article’s critical caveat: “It is unclear if the drone may have been operated in this [an AI-enabled autonomous] way in Ukraine.” Other outlets re-reported the WIRED story, but irresponsibly did so without the caveat. WIRED’s assessment that Kalashnikov claims the KUB-BLA “boasts the ability to identify targets using artificial intelligence” is based on two main pieces of evidence: a Kalashnikov press release about ZALA Aero’s “Artificial Intelligence Visual Identification (AIVI)” capabilities for its unmanned aircraft, and the original Kalashnikov press release announcing the KUB-BLA in 2019. However, these two pieces of evidence are less than they seem. The Russian-language AIVI press release never mentions the KUB-BLA or military applications. Instead, it describes a ZALA Aero machine-learning AI drone product line that is marketed to industrial and agricultural sectors. Incorporating modern machine-learning AI into military applications is significantly more difficult than in industrial or agricultural applications. Modern machine-learning AI using deep neural networks offers the opportunity for incredible gains in performance, but that performance depends on having lots of training data during development. Moreover, that training data needs to closely resemble operational conditions. In general, it is much easier to get such training data from commercial customers than from an enemy military, especially if friendly weapons systems and sensors do not often come within range of enemy ones. The most mature military AI applications are ones like satellite reconnaissance: even in peacetime, satellites get to take a lot of pictures of Russian and Chinese military forces, and those pictures can be digitally labeled by human experts to turn them into training data. Training data is what machine learning AI systems learn from. The combination of a learning algorithm and training data is how AI systems learn to recognize what is in the image. But training data is generally application-specific. Satellite image recognition training data only helps build satellite image recognition AI. One cannot magically use labeled satellite image data to train an AI for a robotic drone’s targeting computer (at least not with today’s technology).



Allen, 5/26 -- Allen, G. C. (2022, May 26). *Russia probably has not used AI-enabled weapons in Ukraine, but that could change*. Russia Probably Has Not Used AI-Enabled Weapons in Ukraine, but That Could Change | Center for Strategic and International Studies. Retrieved June 13, 2022, from <https://www.csis.org/analysis/russia-probably-has-not-used-ai-enabled-weapons-ukraine-could-change>

Getting enough of the right sort of training data to incorporate modern AI into, say, a robotic tank's targeting computer, is a much tougher technical challenge. It is not impossible in principle, but in practice, there are far fewer opportunities to collect the right sort of training data. This is not to say Russia has not tried. In the past, Rostec and Kalashnikov executives have not been shy about their attempts to develop weapons that successfully combine modern AI and combat autonomy, so it would be odd if they had succeeded in doing so with the KUB-BLA and not disclosed it in their marketing materials. Kalashnikov has been heavily promoting the KUB-BLA for both Russian and international customers. What does Kalashnikov say about the KUB-BLA specifically? The 2019 KUB-BLA announcement states that the system has two means of delivering the drone and its explosive warhead to target coordinates: "The target coordinates are specified manually or acquired from [the sensor] payload targeting image." The vague latter description is what led many to assume KUB-BLA was using AI. However, "payload targeting image" is consistent with how many other precision-guided munitions and drone loitering munitions work, including ones that do not use any advanced AI capabilities. A Rostec executive specifically described the KUB-BLA as a Russian "domestic analogue" to the Israeli-built Orbiter 1K drone, which looks nearly identical. The Orbiter 1K comes with a ground control station where human operators monitor the video coming from the drone's sensor and select targets directly from the video feed. In other words, a human has already selected the target prior to the drone attacking it, and the drone is only autonomously maintaining target lock and navigating to the target, not autonomously selecting and deciding to engage targets. Autonomy over decisions to "select and engage targets" is the specific standard in U.S. Department of Defense policy as to what qualifies as an "autonomous weapon system." Fire and forget munitions—which is the standard term used to describe not only the Orbiter 1K, but also heat-seeking missiles like the Javelin and Stinger—do not qualify as autonomous weapons. Those heat-seeking missiles do not use modern deep neural network machine learning, but they do use thermal image processing algorithms that were once considered state of the art. They, and many more systems like them, have been in use for decades by dozens of militaries around the world. It is possible that the KUB-BLA received some kind of lethal AI-targeting upgrade prior to being used in Ukraine, but that is doubtful. Neither WIRED nor anyone else has provided evidence that this is the case. If the weapon did have those capabilities, it is unlikely that Kalashnikov would fail to mention them. The company has bragged about KUB-BLA's recent use in Ukraine, and in the past Kalashnikov has openly talked about seeking to develop a "fully automated combat module" based on AI deep neural network technology.



National Association for Urban Debate Leagues

Responds to Laird 20 – cooperation with AI fosters diplomacy

EVEN WITH LIMITED ADVANCEMENT, INCREASED ACCESSIBILITY EMBOLDENS RUSSIA TO DEPLOY LAWS WITHOUT DATA TRAINING

Allen, 5/26 -- Allen, G. C. (2022, May 26). *Russia probably has not used AI-enabled weapons in Ukraine, but that could change.* Russia Probably Has Not Used AI-Enabled Weapons in Ukraine, but That Could Change | Center for Strategic and International Studies. Retrieved June 13, 2022, from <https://www.csis.org/analysis/russia-probably-has-not-used-ai-enabled-weapons-ukraine-could-change>

In sum, there is little reason to believe that Russia is using AI-enabled autonomous weapons in Ukraine, yet. That is the good news. **The bad news is that, if Russia's unlawful war in Ukraine drags on, Russia has the intent and likely has the means to deploy autonomous weapons, with or without advanced AI.** Regarding means, a recent report by Russian news outlet RIA Novosti interviewed an unnamed Russian military source that is worth quoting (via Microsoft's automatic translation) at length: Russian reconnaissance and reconnaissance-strike UAVs will receive a digital catalog with electronic [optical and infrared] images of military equipment adopted in NATO countries. This will allow them to automatically identify it on the battlefield and create a map of the location of enemy positions directly onboard the device, which will be broadcast to the command post. . . . It is formed due to neural network training algorithms, which makes it possible to accurately determine the samples of equipment in a wide variety of environmental conditions, including with a short exposure (the technique is visible for several seconds or less), as well as when only part of the sample falls into the field of view of the drone—when, for example, only part of any combat vehicle is visible from cover. As mentioned above, **collecting adequate training data remains a significant hurdle for many military AI development projects.** While the invasion of Ukraine has been a disaster in many ways for the Russian military, **NATO has provided weapons and equipment to Ukraine that offers the best opportunity yet to collect operational training data for new AI models and more diverse military AI applications. The anonymous quote suggests that Russia's military is taking this opportunity seriously.** Of course, domestic opposition to the war has caused an exodus of tech workers from Russia, and the sanctions levied against Russia have left it with major shortages of the semiconductor chips needed to make advanced AI systems. These are major challenges, but **advanced AI is not required to endow weapons with all types of lethal autonomous capabilities, only a willingness to delegate decisions and freedom to military machines.** The Israeli-built Harpy autonomous weapon, which can loiter in the air over a battlefield for hours while searching for enemy radar emissions to attack, dates back to the late 1980s. In addition to the KUB-BLA, Kalashnikov makes another drone called the Lancet, which the aforementioned Rostec executive describes as a Russian analogue to the more modern Harpy-2 (aka Harop). Kalashnikov claims the following about the Lancet: [The Lancet] is a smart multipurpose weapon, capable of autonomously finding and hitting a target. The weapon system consists of precision strike component, reconnaissance, navigation and communications modules. It creates its own navigation field and does not require ground or sea-based infrastructure.



Allen, 5/26 -- Allen, G. C. (2022, May 26). *Russia probably has not used AI-enabled weapons in Ukraine, but that could change*. Russia Probably Has Not Used AI-Enabled Weapons in Ukraine, but That Could Change | Center for Strategic and International Studies. Retrieved June 13, 2022, from <https://www.csis.org/analysis/russia-probably-has-not-used-ai-enabled-weapons-ukraine-could-change>

Kalashnikov is clearly marketing the Lancet as a capable autonomous weapon but also one that can be remotely controlled depending on user preference. The Russian military has already used the Lancet for combat operations in Syria in its remotely controlled mode, but observers have not yet confirmed that the Lancet is being used in Ukraine. Once it shows up, Russia will likely be tempted to turn on Lancet's autonomous weapon functionality—that is, if the system's performance matches Kalashnikov's advertising. Remotely piloted drones have demonstrated effectiveness in the war in Ukraine that exceeded most analysts' expectations before the war, when drones were often viewed as useful in counterinsurgency operations, but not a high-end conflict against a technologically sophisticated adversary like Russia. In military technology competition, however, each successful move leads to a countermove. Many analysts feel that the weak point in current drone weapons is their reliance on a high-bandwidth communications link to their human remote controllers. If the war in Ukraine drags on for many more months or years, expect both sides to more widely deploy jammers and other electronic warfare systems to counter drones. Elimination or reduction of remote piloting options will naturally lead Russia to desire greater autonomy in their drone weapons. Regarding intent, Russia has consistently played the role of stubborn obstructionist through years of UN expert discussions about developing new international norms or codes of conduct for the development and use of autonomous weapons. Most countries around the world are being cautious with the introduction of military AI. Much of the AI ethics and autonomous weapons debate has been heavily focused on whether or not using AI increases the risk of technical accidents and harm to civilians. But Russia's unprovoked war in Ukraine is a tragic reminder that, while unintentional harm to civilians is a real tragedy, there is also the unsolved problem of intentional harm to civilians. The Russian military has routinely attacked not only residential neighborhoods, but hospitals and humanitarian organizations. Finally, Russia's human soldiers in Ukraine have suffered heavy losses and reportedly deserted in large numbers. Faced with such frustrations, there's little reason to doubt Russian president Vladimir Putin would use lethal autonomous weapons if he thought it would provide a military edge. The United States and its allies need to start thinking about how to ensure that he does not.



SOLVENCY

Responds to Franke 21 – AI coop maintains US tech leadership

NATO MUST SHIFT FOCUS NOW. TECHNOLOGICAL ADVANCES CANNOT SUCCEED WITHOUT PERSONNEL TRAINING AND ADAPTATION

Dolan, 6/8 -- Dolan, C. (2022, June 8). NATO's 2022 strategic concept must enhance digital access and capacities. Just Security. Retrieved June 12, 2022, from <https://www.justsecurity.org/81839/natos-2022-strategic-concept-must-enhance-digital-access-and-capacities/>

To deter hybrid threats across multiple domains, with enhanced access on different digital platforms, NATO members should develop smarter and lethal capabilities to confront threats from state and non-state actors. This would allow ACT and ACO to prepare for any contingency and respond to adversaries in battlefields and battlespaces. Plug and Play The 2022 Strategic Concept encourages collaboration in the implementation of guidelines and procedures through a “plug-and-play” concept, in which platforms and systems are optimized for readiness and response at lightning speed. Plug-and-play is based on approaches used in commercial software that allow for innovation and easy access to networks and systems through secure platforms. The NATO School Oberammergau should offer platform training and education courses programs in mobile access for ACT and ACO personnel with appropriate security clearances. This would allow them to access the appropriate platform and utilize data and information necessary for their tasks and responsibilities. For example, NATO’s Small Arms and Light Weapons (SALW) and Mine Action (MA) Information Sharing Platform contains rich and publicly available datasets on the roles played by the alliance in mitigating the illicit trade in small arms, tanks, aircraft, and naval vessels. It reports and updates NATO-funded projects to prevent adversaries from acquiring these weapons. However, the SALW-MA platform is outdated and not user friendly, impeding its functionality in practice. Put simply, NATO’s ACT and ACO should focus as much on easing access to information as it does on advanced technologies and conventional weaponry. This would provide NATO with useful tools to access data and intelligence on the strategic, operational, and tactical levels and in land, sea, air, space, and cyber domains using devices and platforms that can seamlessly connect in different locations. But NATO Commands cannot simply expect its existing personnel to adapt. They must be trained and educated on a regular basis to use digital infrastructures in ways that make their jobs easier. On the strategic level, the 2022 Strategic Concept must provide NATO’s political and military leaders with flexibility and resources to discern the diversity of hybrid threats in the environment. NATO’s strategic planners, cyber operators, and warfighters should be trained and educated in relevant digital platforms, access, and sharing data and information in ways that improve collaborative decision-making and collective defense. On the operations level, personnel must be given enough space to share data and intelligence as well as to train tactical level personnel on software that enables them to collect, analyze, process, and disseminate information quickly and easily across multiple domains. Addressing these challenges is difficult for just one nation-state, let alone for all 30 NATO members. Therefore, the 2022 Strategic Concept should emphasize connectivity between member states in multi-domain operations and in collaboration with the private sector and academia. Accessibility to information and data sharing among NATO members should be securitized and harmonized.



AI IS DETRIMENTAL TO HUMAN RIGHTS AND PRIVACY. THROUGH THE PRIVATE SECTOR AND GOVERNMENTS AROUND THE GLOBE, WE ARE AT AN UNPRECEDENTED LEVEL OF SURVEILLANCE

Dziedzic, 21 -- Dziedzic, M. (2021, September 15). *Urgent action needed over artificial intelligence risks to human rights* | | UN news. United Nations. Retrieved June 14, 2022, from <https://news.un.org/en/story/2021/09/1099972>

Urgent action is needed as it can take time to assess and address the serious risks this technology poses to human rights, warned the High Commissioner: “The higher the risk for human rights, the stricter the legal requirements for the use of AI technology should be”. Ms. Bachelet also called for AI applications that cannot be used in compliance with international human rights law, to be banned. “Artificial intelligence can be a force for good, helping societies overcome some of the great challenges of our times. But AI technologies can have negative, even catastrophic, effects if they are used without sufficient regard to how they affect people’s human rights”. Pegasus spyware revelations On Tuesday, the UN rights chief expressed concern about the “unprecedented level of surveillance across the globe by state and private actors”, which she insisted was “incompatible” with human rights. She was speaking at a Council of Europe hearing on the implications stemming from July’s controversy over Pegasus spyware. The Pegasus revelations were no surprise to many people, Ms. Bachelet told the Council of Europe’s Committee on Legal Affairs and Human Rights, in reference to the widespread use of spyware commercialized by the NSO group, which affected thousands of people in 45 countries across four continents. ‘High price’, without action The High Commissioner’s call came as her office, OHCHR, published a report that analyses how AI affects people’s right to privacy and other rights, including the rights to health, education, freedom of movement, freedom of peaceful assembly and association, and freedom of expression. The document includes an assessment of profiling, automated decision-making and other machine-learning technologies. The situation is “dire” said Tim Engelhardt, Human Rights Officer, Rule of Law and Democracy Section, who was speaking at the launch of the report in Geneva on Wednesday. The situation has “not improved over the years but has become worse” he said. Whilst welcoming “the European Union’s agreement to strengthen the rules on control” and “the growth of international voluntary commitments and accountability mechanisms”, he warned that “we don’t think we will have a solution in the coming year, but the first steps need to be taken now or many people in the world will pay a high price”. OHCHR Director of Thematic Engagement, Peggy Hicks, added to Mr Engelhardt’s warning, stating “it’s not about the risks in future, but the reality today. Without far-reaching shifts, the harms will multiply with scale and speed and we won’t know the extent of the problem.”

.....



Gubrud, 3/2 – Gubrud, M. (2022, March 2). *Why should we ban autonomous weapons? to survive*. IEEE Spectrum. Retrieved June 8, 2022, from <https://spectrum.ieee.org/why-should-we-ban-autonomous-weapons-to-survive>

-- wasn't about banning teeny-tiny Kalashnikovs, but identifying the qualitatively distinct new things that emerging technologies would enable. One of my first ideas was a ban on autonomous kill decision by machines. "I knew that most people would agree we should not have killer robots, but when I started talking about banning them, people would mostly stare" I knew that most people would agree we should not have killer robots. This made lethal autonomy a bright red line at which it might be possible to erect a roadblock to the arms race. I also knew that unless we resolved not to cross that line, we would soon enter an era in which, once the fighting had started, the complexity and speed of automated combat, and the delegation of lethal autonomy as a military necessity, would put the war machines effectively beyond human control. But when I started to talk about banning killer robots, people would mostly stare. Military people angrily denied that anyone would even consider letting machines decide when to fire guns and at what or at whom. For many years the U.S. military resisted autonomous weapons, concerned about their legality, controllability and potential for friendly-fire accidents. Systems like the CAPTOR mine, designed to autonomously launch a homing torpedo at a passing submarine, and the LOCAAS mini-cruise missile, designed to loiter above a battlefield and search for tanks or people to kill, were canceled or phased out. As late as 2013, a poll conducted by Charli Carpenter, a political science professor at the University of Massachusetts Amherst, found Americans against using autonomous weapons by 2-to-1, and tellingly, military personnel were among those most opposed to killer robots. Yet starting in 2001, the use of armed drones by the United States began to make the question of future autonomous weapons more urgent. In a 2004 article, Juergen Altmann and I declared that "Autonomous 'killer robots' should be prohibited" and added that "a human should be the decision maker when a target is to be attacked." In 2009, Altmann, a professor of physics at Technische Universität Dortmund, co-founded the International Committee for Robot Arms Control, and at its first conference a year later, I suggested human control as a fundamental principle. The unacceptability of machine decision in the use of violent force could be asserted, I argued, without need of scientific or legal justification. In 2012, Human Rights Watch began to organize the Campaign to Stop Killer Robots, a global coalition that now includes more than 60 nongovernmental organizations. The issue rose to prominence with astonishing speed, and the United Nations Convention on Certain Conventional Weapons (CCW) held its first "Meeting of Experts on Lethal Autonomous Weapon Systems" in May 2014, and another the following year. This past April, the third such meeting concluded with a recommendation to form a "Group of Governmental Experts," the next step in the process of negotiating... something. Many statements at the CCW have endorsed human control as a guiding principle, and Altmann and I have suggested cryptographic proof of accountable human control as a way to verify compliance with a ban on autonomous weapons. Yet the CCW has not set a definite goal for its deliberations. And in the meantime, the killer robot arms race has taken off. FULL SPEED AHEAD In 2012, the Obama administration, via then-undersecretary of defense Ashton Carter, directed the Pentagon to begin developing, acquiring, and using "autonomous and semi-autonomous weapon systems." Directive 3000.09 has been widely misperceived as a policy of caution; many accounts insist that it "requires a human in the loop." But instead of human control, the policy sets "appropriate levels of human judgment" as a guiding principle. It does not explain what that means, but senior officials are required --



AT: "LAWS ARE GOOD FOR MILITARY OPERATIONS"

OVER INVESTMENT IN AI CAN LEAD TO OVERDEPENDENCE, LEADING TO FATALLY UNINTENTIONAL CONSEQUENCES SUCH AS NUCLEAR WAR.

Kallenborn, 22 (Zachary Kallenborn, Zachary Kallenborn is a research affiliate with the Unconventional Weapons and Technology Division of the National Consortium for the Study of Terrorism and Responses to Terrorism (START), a policy fellow at the Schar School of Policy and Government, a US Army Training and Doctrine Command "Mad Scientist," and national security consultant, 2-1-2022, accessed on 6-13-2022, Bulletin of the Atomic Scientists, "Giving an AI control of nuclear weapons: What could possibly go wrong? - Bulletin of the Atomic Scientists", <https://thebulletin.org/2022/02/giving-an-ai-control-of-nuclear-weapons-what-could-possibly-go-wrong/>)

How autonomous nuclear weapons could go wrong. The huge problem with autonomous nuclear weapons, and really all autonomous weapons, is error. Machine learning-based artificial intelligences—the current AI vogue—rely on large amounts of data to perform a task. Google's AlphaGo program beat the world's greatest human go players, experts at the ancient Chinese game that's even more complex than chess, by playing millions of games against itself to learn the game. For a constrained game like Go, that worked well. But in the real world, data may be biased or incomplete in all sorts of ways. For example, one hiring algorithm concluded being named Jared and playing high school lacrosse was the most reliable indicator of job performance, probably because it picked up on human biases in the data.

In a nuclear weapons context, a government may have little data about adversary military platforms; existing data may be structurally biased, by, for example, relying on satellite imagery; or data may not account for obvious, expected variations such as imagery in taken during foggy, rainy, or overcast weather.

AND, WARS UTILIZING AUTONOMOUS WEAPONS ARE MORE LIKELY TO ESCALATE DUE TO PREEMPTIVE ATTACKS.

Laird, 16 (Burgess Laird, Burgess Laird is a Senior International Defense Researcher with the RAND Corporation and an adjunct instructor in the M.A. in Global Security Studies at Johns Hopkins University, 12-8-2016, accessed on 6-9-2022, Rand, "The Risks of Autonomous Weapons Systems for Crisis Stability and Conflict Escalation in Future U.S.-Russia Confrontations", <https://www.rand.org/blog/2020/06/the-risks-of-autonomous-weapons-systems-for-crisis.html>)

First, a state facing an adversary with AWS capable of making decisions at machine speeds is likely to fear the threat of sudden and potent attack, a threat that would compress the amount of time for strategic decisionmaking. The posturing of AWS during a crisis would likely create fears that one's forces could suffer significant, if not decisive, strikes. These fears in turn could translate into pressures to strike first—to preempt—for fear of having to strike second from a greatly weakened position. Similarly, within conflict, the fear of losing at machine speeds would be likely to cause a state to escalate the intensity of the conflict possibly even to the level of nuclear use.



Gubrud, 3/2 – Gubrud, M. (2022, March 2). *Why should we ban autonomous weapons? to survive*. IEEE Spectrum. Retrieved June 8, 2022, from <https://spectrum.ieee.org/why-should-we-ban-autonomous-weapons-to-survive>

Killer robots pose a threat to all of us. In the movies, this threat is usually personified as an evil machine bent on destroying humanity for reasons of its own. In reality, the threat comes from within us. It is the threat of war. In today's drone warfare, people kill other people from the safety of cubicles far away. Many do see something horrific in this. Even more are horrified by the idea of replacing the operator with artificial intelligence, and dispatching autonomous weapons to hunt and kill without further human involvement. Proponents of autonomous weapons say their use is inevitable and natural, a mere extension of human will and judgment through the agency of machines. They question whether artificial intelligence will always be incapable of distinguishing civilians from combatants, or even of making reasonable tradeoffs between military gains and risk or harm to civilians. After all, they argue, people are often cruel and stupid, and soldiers under extreme stress sometimes go berserk and commit atrocities. What if autonomous weapons, used judiciously, could actually save lives of soldiers and civilians? I'll agree that we can imagine circumstances in which using an intelligent autonomous weapon could cause less harm than a more destructive, dumb weapon, if those were the only choices. But human-controlled robotic weapons could often be just as effective, or it might be possible to avoid violence altogether. Autonomous weapons could malfunction, kill innocents, and nobody be held responsible. Which kind of situation would occur most often, and whether autonomous weapons would be more or less deadly than their prohibition, assuming everything else would be the same, is endlessly debatable. "The major powers are developing autonomous missiles and drones that will hunt ships, subs, and tanks, and piecing together highly automated battle networks that will confront each other and have the capability of operating without human control" But everything else won't be the same. Proponents claim that machine intelligence and autonomous weapons will revolutionize warfare, and that no nation can risk letting its enemies have a monopoly on them. Even if this is exaggerated, it shows the potential for a strong stimulus to the global arms race. These technologies are being pursued most vigorously by the nuclear-armed nations. In the United States, they are touted as the answer to rising challenges from China and Russia, as well as from lesser powers armed with modern weaponry. The major powers are developing autonomous missiles and drones that will hunt ships, subs, and tanks, and piecing together highly automated battle networks that will confront each other and have the capability of operating without human control. **Autonomous weapons are a salient point of departure in a technology-fueled arms race that puts everyone in danger. That is why I believe we need to ban them as fast and as hard as we possibly can.** A BRIGHT RED LINE It's a view I've held for almost three decades, and it wasn't inspired by the Terminator, but by the 1988 incident in which a U.S. Navy air defense system mistakenly shot down an Iranian airliner. Although human error appears to have played the deciding role in that incident, part of the problem was excessive reliance on complex automated systems under time pressure and uncertain warnings of imminent danger—the classic paradigm for "accidental war." At the time, as an intern at the Federation of American Scientists in Washington, D.C., I was looking at nanotechnology and the rush of new capabilities that would come as we learn to build ever more complex systems with ever smaller parts. We see that today in billion-transistor chips and the computers, robots, and machine learning systems they are making possible. I worried about a runaway arms race. I was asked to come up with proposals for nanotechnology arms control. I decided it --



NATIONS AND THEIR MILITARIES MIRROR US – ADVANCING AI TECH NOW. US PRODUCTION IS PAVING WAY FOR SHORTCUTS AND EARLY USE OF AI

Gubrud, 3/2 – Gubrud, M. (2022, March 2). *Why should we ban autonomous weapons? to survive*. IEEE Spectrum. Retrieved June 8, 2022, from <https://spectrum.ieee.org/why-should-we-ban-autonomous-weapons-to-survive>

-- to certify that autonomous weapon systems meet this standard if they select and kill people without human intervention. The policy clearly does not forbid such systems. Rather, it permits the withdrawal of human judgment, control, and responsibility from points of lethal decision. The policy has not stood in the way of programs such as the Long Range Anti-Ship Missile, slated for deployment in 2018, which will hunt its targets over a wide expanse, relying on its own computers to discriminate enemy ships from civilian vessels. Weapons like this are classified as merely "semi-autonomous" and get a green light without certification, even though they will be operating fully autonomously when they decide which pixels and signals correspond to valid targets, and attack them with lethal force. Every technology needed to acquire, track, identify, and home in or control firing on targets can be developed and used in "semi-autonomous weapon systems," which can even be sent on hunt-and-kill missions as long as the quarry has been "selected by a human operator." (In case you're wondering, "target selection" is defined as "The determination that an individual target or a specific group of targets is to be engaged.") It's unclear that the policy stands in the way of anything. In reality, the directive signaled an upward inflection in the trend toward killer robots. Throughout the military there is now open discussion about autonomy in future weapon systems; ambitious junior officers are tying their careers to it. DARPA and the Navy are particularly active in efforts to develop autonomous systems, but the Air Force, Army, and Marines won't be left out. Carter, now the defense secretary, is heavily promoting AI and robotics programs, establishing an office in Silicon Valley and a board of advisors to be chaired by Eric Schmidt, the executive chairman of Google's parent company Alphabet. The message has been received globally as well. Russia in 2013 moved to create its own versions of DARPA and the of the U.S. Navy's Laboratory for Autonomous Systems Research, and deputy prime minister Dmitry Rogozin called on Russian industry to create weapons that "strike on their own," pointing explicitly to American developments. China, too, has been developing its own drones and robotic weapons, mirroring the United States (but with less noise than Russia). Britain, Israel, India, South Korea... in fact, every significant military power on Earth is looking in this direction. "The United States has been leading the robot arms race, both with weapons development and with a policy that pretends to be cautious and responsible but actually clears the way for vigorous development and early use of autonomous weapons" Both Russia and China have engaged in aggressive actions, arms buildups, and belligerent rhetoric in recent years, and it's unclear whether they could be persuaded to support a ban on autonomous weapons. But we aren't even trying. Instead, the United States has been leading the robot arms race, both with weapons development and with a policy that pretends to be cautious and responsible but actually clears the way for vigorous development and early use of autonomous weapons. Deputy defense secretary Robert Work has championed the notion of a "Third Offset" in which the United States would leap to the next generation of military technologies ahead of its "adversaries," particularly Russia and China. To calm fears about robots taking over, he emphasizes "human-machine collaboration and combat teaming" and says the military will use artificial intelligence and robotics to augment, not replace human warfighters. Yet he worries that adversaries may field fully autonomous weapon systems, and says the U.S. may need to "delegate authority to machines" because "humans simply cannot operate at the same speed."



PROLIFERATION OF LAWS CAUSES INTERNATIONAL POWER WAR. THIS LEADS TO CONFLICT BETWEEN RUSSIA, CHINA, AND THE US TO LEAD THE ARMS RACE

Gubrud, 3/22 – Gubrud, M. (2022, March 2). *Why should we ban autonomous weapons? to survive*. IEEE Spectrum. Retrieved June 8, 2022, from <https://spectrum.ieee.org/why-should-we-ban-autonomous-weapons-to-survive>

Work admits that the United States has no monopoly on the basic enabler, information technology, which today is driven more by commercial markets than by military needs. Both China and Russia have strong software and cyber hacking capabilities. Their latest advanced fighters, tanks, and missiles are said to rival ours in sophistication. Work compares the present to the “inter-war period” and urges the U.S. to emulate Germany’s invention of blitzkrieg. Has he forgotten how that ended? **Nobody wants war. Yet, fearing enemy aggression, we position ourselves at the brink of it.** Arms races militarize societies, inflate threat perceptions, and yield a proliferation of opportunities for accidents and mistakes. In numerous close calls during the Cold War, it came down to the judgment of one or a few people not to take the next step in a potentially fatal chain of events. But machines simply execute their programs, as intended. They also behave in ways we did not intend or expect. “Networks of autonomous weapons could accidentally ignite a war and, once it has started, rapidly escalate it out of control. To set up such a disaster waiting to happen would be foolish” Our experience with the unpredictable failures and unintended interactions of complex software systems, particularly competitive autonomous agents designed in secrecy by hostile teams, serves as a warning that networks of autonomous weapons could accidentally ignite a war and, once it has started, rapidly escalate it out of control. To set up such a disaster waiting to happen would be foolish, but not unprecedented. It’s the type of risk we took during the Cold War, and it’s similar to the military planning that drove the march to war in 1914. Arms races and confrontation push us to take this kind of risk. Paul Scharre, one of the architects of Directive 3000.09, has suggested that the risk of autonomous systems acting on their own could be mitigated by negotiating “rules of the road” and including humans in battle networks as “fail-safes.” But it’s asking a lot of humans to remain calm when machines indicate an attack underway. By the time you sort out a false alarm, autonomous weapons may actually have started fighting. If nations can’t agree to the simple idea of a verified ban to avoid this danger, it seems less likely that they will be able to negotiate some complicated system of rules and safeguards. Direct authority to launch a nuclear strike may never be delegated to machines and a war between the United States and China or Russia might not end in nuclear war, but do we want to take that risk? There is no reason to believe we can engineer safety into a tense confrontation between networks of autonomous weapons at the brink of war. The further we go down that road, the harder it will be to walk back. Banning autonomous weapons and asserting the primacy of human control isn’t a complete solution, but it is probably an essential step to ending the arms race and building true peace and security. The fundamental problem is conflict itself, which pits human against human, reason against reason and machine against machine. We struggle to contain our conflicts, but passing them on to machines risks finding ourselves nominally still in command yet unable to control events at superhuman speed. We are horrified by killer robots, and we can ground their prohibition on strong a priori principles such as human control, responsibility, dignity—and survival. Instead of endlessly debating the validity of these human prejudices, we should take them as saving grace, and use them to stop killer robots.



AI RESEARCHERS AND DEVELOPERS CALL A STOP TO PRODUCTION. AI ADVANCEMENT LEADS TO LAUNDRY LIST OF “BAD ACTOR” SCENARIOS

Conn, 18 -- Conn, A. (2018, September 4). *The risks posed by Lethal Autonomous Weapons*. Future of Life Institute. Retrieved June 8, 2022, from <https://futureoflife.org/2018/09/04/the-risks-posed-by-lethal-autonomousweapons/?cn-reloaded=1>

Killer robots. It's a phrase that's both terrifying, but also one that most people think of as still in the realm of science fiction. Yet weapons built with artificial intelligence (AI) – weapons that could identify, target, and kill a person all on their own – are quickly moving from sci-fi to reality. To date, no weapons exist that can specifically target people. But there are weapons that can track incoming missiles or locate enemy radar signals, and these weapons can autonomously strike these non-human threats without any person involved in the final decision. Experts predict that in just a few years, if not sooner, this technology will be advanced enough to use against people. Over the last few years, delegates at the United Nations have debated whether to consider banning killer robots, more formally known as lethal autonomous weapons systems (LAWS). This week delegates met again to consider whether more meetings next year could lead to something more tangible – a political declaration or an outright ban. Meanwhile, those who would actually be responsible for designing LAWS — the AI and robotics researchers and developers — have spent these years calling on the UN to negotiate a treaty banning LAWS. More specifically, nearly 4,000 AI and robotics researchers called for a ban on LAWS in 2015; in 2017, 137 CEOs of AI companies asked the UN to ban LAWS; and in 2018, 240 AI-related organizations and nearly 3,100 individuals took that call a step further and pledged not to be involved in LAWS development. And AI researchers have plenty of reasons for their consensus that the world should seek a ban on lethal autonomous weapons. Principle among these is that AI experts tend to recognize how dangerous and destabilizing these weapons could be. The weapons could be hacked. The weapons could fall into the hands of “bad actors.” The weapons may not be as “smart” as we think and could unwittingly target innocent civilians. Because the materials necessary to build the weapons are cheap and easy to obtain, military powers could mass-produce these weapons, increasing the likelihood of proliferation and mass killings. The weapons could enable assassinations or, alternatively, they could become weapons of oppression, allowing dictators and warlords to subdue their people.



VIOLENCE FOLLOWS

Conn, 18 -- Conn, A. (2018, September 4). *The risks posed by Lethal Autonomous Weapons*. Future of Life Institute. Retrieved June 8, 2022, from <https://futureoflife.org/2018/09/04/the-risks-posed-by-lethal-autonomousweapons/?cn-reloaded=1>

But perhaps the greatest risk posed by LAWS, is the potential to ignite a global AI arms race. For now, governments insist they will ensure that testing, validation, and verification of these weapons is mandatory. However, these weapons are not only technologically novel, but also transformative; they have been described as the third revolution in warfare, following gun powder and nuclear weapons. LAWS have the potential to become the most powerful types of weapons the world has seen. Varying degrees of autonomy already exist in weapon systems around the world, and levels of autonomy and advanced AI capabilities in weapons are increasing rapidly. If one country were to begin substantial development of a LAWS program — or even if the program is simply perceived as substantial by other countries — an AI arms race would likely be imminent. During an arms race, countries and AI labs will feel increasing pressure to find shortcuts around safety precautions. Once that happens, every threat mentioned above becomes even more likely, if not inevitable. As stated in the Open Letter Against Lethal Autonomous Weapons: The key question for humanity today is whether to start a global AI arms race or to prevent it from starting. If any major military power pushes ahead with AI weapon development, a global arms race is virtually inevitable, and the endpoint of this technological trajectory is obvious: autonomous weapons will become the Kalashnikovs of tomorrow. Unlike nuclear weapons, they require no costly or hard-to-obtain raw materials, so they will become ubiquitous and cheap for all significant military powers to mass-produce. It will only be a matter of time until they appear on the black market and in the hands of terrorists, dictators wishing to better control their populace, warlords wishing to perpetrate ethnic cleansing, etc. Autonomous weapons are ideal for tasks such as assassinations, destabilizing nations, subduing populations and selectively killing a particular ethnic group. Most countries here have expressed their strong desire to move from talking about this topic to reaching an outcome. There have been many calls from countries and groups of countries to negotiate a new treaty to either prohibit LAWS and/or affirm meaningful human control over the weapons. Some countries have suggested other measures such as a political declaration. But a few countries — especially Russia, the United States, South Korea, Israel, and Australia — are obfuscating the process, which could lead us closer to an arms race. This is a threat we must prevent.



REGULATIONS EXT.

US HAS A HISTORY OF MISUSING AI AND SPURRING MISTRUST IN THE TECHNOLOGY, FLINT PROVES

Hendler, 19 -- Hendler, J. (2019, April 18). *As governments adopt artificial intelligence, there's little oversight and lots of Danger*. The Conversation. Retrieved June 14, 2022, from <https://theconversation.com/as-governments-adopt-artificial-intelligence-theres-little-oversight-and-lots-of-danger-115109>

Artificial intelligence systems can – if properly used – help make government more effective and responsive, improving the lives of citizens. Improperly used, however, the dystopian visions of George Orwell's "1984" become more realistic. On their own and urged by a new presidential executive order, governments across the U.S., including state and federal agencies, are exploring ways to use AI technologies. As an AI researcher for more than 40 years, who has been a consultant or participant in many government projects, I believe it's worth noting that sometimes they've done it well – and other times not quite so well. The potential harms and benefits are significant. An early success In 2015, the U.S. Department of Homeland Security developed an AI system called "Emma," a chatbot that can answer questions posed to it in regular English, without needing to know what "her" introductory website calls "government speak" – all the official terms and acronyms used in agency documents. By late 2016, DHS reported that Emma was already helping to answer nearly a half-million questions per month, allowing DHS to handle many more inquiries than it had previously, and letting human employees spend more time helping people with more complicated queries that are beyond Emma's abilities. This sort of conversation-automating artificial intelligence has now been used by other government agencies, in cities and countries around the world. Flint's water A more complicated example of how governments could aptly apply AI can be seen in Flint, Michigan. As the local and state governments struggled to combat lead contamination in the city's drinking water, it became clear that they would need to replace the city's remaining lead water pipes. However, the city's records were incomplete, and it was going to be extremely expensive to dig up all the city's pipes to see if they were lead or copper. For a time, artificial intelligence analysis helped guide pipe replacement in Flint, Michigan. Instead, computer scientists and government employees collaborated to analyze a wide range of data about each of 55,000 properties in the city, including how old the home was, to calculate the likelihood it was served by lead pipes. Before the system was used, 80% of the pipes dug up needed to be replaced, which meant 20% of the time, money and effort was being wasted on pipes that didn't need replacing. The AI system helped engineers focus on high-risk properties, identifying a set of properties most likely to need pipe replacements. When city inspectors visited to verify the situation, the algorithm was right 70% of the time. That promised to save enormous amounts of money and speed up the pipe replacement process. However, local politics got in the way. Many members of the public didn't understand why the system was identifying the homes it did, and objected, saying the AI method was unfairly ignoring their homes. After city officials stopped using the algorithm, only 15% of the pipes dug up were lead. That made the replacement project slower and more costly. Distressing examples The problem in Flint was that people didn't understand that AI technology was being used well, and that people were verifying its findings with independent inspections. In part, this was because they didn't trust AI – and in some cases there is good reason for that.



EXPERTS AGREE – AI IS BIASED, UNRELIABLE, INACCURATE, AND PAVES THE WAY FOR GOVERNMENTS TO MISUSE THE TECH IN THE FUTURE

Hendler, 19 -- Hendler, J. (2019, April 18). *As governments adopt artificial intelligence, there's little oversight and lots of Danger*. The Conversation. Retrieved June 14, 2022, from <https://theconversation.com/as-governments-adopt-artificial-intelligence-theres-little-oversight-and-lots-of-danger-115109>

In 2017, I was among a group of more than four dozen AI researchers who sent a letter to the acting secretary of the U.S. Department of Homeland Security. We expressed concerns about a proposal to use automated systems to determine whether a person seeking asylum in the U.S. would become a “positively contributing member of society” or was more likely to be a terrorist threat. “Simply put,” our letter said, “no computational methods can provide reliable or objective assessments of the traits that [DHS] seeks to measure.” We explained that machine learning is susceptible to a problem called “data skew,” in which the system’s ability to predict a characteristic depends in part on how common that characteristic is in the data used to train the system. A face tracking and analysis system takes a look at a woman’s face. So in a database of 300 million Americans, if one in 100 people are, say, of Indian descent, the system will be fairly accurate at identifying them. But if looking at a characteristic shared by just one in a million Americans, there really isn’t enough data for the algorithm to make a good analysis. As the letter explained, “on the scale of the American population and immigration rates, criminal acts are relatively rare, and terrorist acts are extremely rare.” Algorithmic analysis is extremely unlikely to identify potential terrorists. Fortunately, our arguments proved convincing. In May 2018, DHS announced it would not use a machine learning algorithm in this way. Other worrying efforts Other government uses of AI are being questioned, too – such as attempts at “predictive policing,” setting bail amounts and criminal sentences and hiring government workers. All of these have been shown to be susceptible to technical issues and data limitations that can bias their decisions based on race, gender or cultural background. Other AI technologies such as facial recognition, automated surveillance and mass data collection are raising real concerns about security, privacy, fairness and accuracy in a democratic society. As Trump’s executive order demonstrates, there is significant interest in harnessing AI for its fullest positive potential. But the significant dangers of abuse, misuse and bias – whether intentional or not – have the potential to work against the very principles international democracies have been built upon. As the use of AI technologies grows, whether originally well-meant or deliberately authoritarian, the potential for abuse increases as well. With no currently existing government-wide oversight in place in the U.S., the best way to avoid these abuses is teaching the public about the appropriate uses of AI by way of conversation between scientists, concerned citizens and public administrators to help determine when and where it is inappropriate to deploy these powerful new tools.



Thomas, 21 -- Thomas, M. (2021, July 7). *7 dangerous risks of Artificial Intelligence*. Built In. Retrieved June 8, 2022, from <https://builtin.com/artificial-intelligence/risks-of-artificial-intelligence>

Having exterminated humanity, it then calculates the number with higher confidence." That's jarring, sure. But rest easy. In 2012 the Obama Administration's Department of Defense issued a directive regarding "Autonomy in Weapon Systems" that included this line: "Autonomous and semi-autonomous weapon systems shall be designed to allow commanders and operators to exercise appropriate levels of human judgment over the use of force." And in early November of this year, a Pentagon group called the Defense Innovation Board published ethical guidelines regarding the design and deployment of AI-enabled weapons. According to the Washington Post, however, "the board's recommendations are in no way legally binding." It now falls to the Pentagon to determine how and whether to proceed with them." Well, that's a relief. Or not. Have you ever considered that algorithms could bring down our entire financial system? That's right, Wall Street. You might want to take notice. Algorithmic trading could be responsible for our next major financial crisis in the markets. What is algorithmic trading? This type of trading occurs when a computer, unencumbered by the instincts or emotions that could cloud a human's judgement, execute trades based off of pre-programmed instructions. These computers can make extremely high-volume, high-frequency and high-value trades that can lead to big losses and extreme market volatility. Algorithmic High-Frequency Trading (HFT) is proving to be a huge risk factor in our markets. HFT is essentially when a computer places thousands of trades at blistering speeds with the goal of selling a few seconds later for small profits. Thousands of these trades every second can equal a pretty hefty chunk of change. The issue with HFT is that it doesn't take into account how interconnected the markets are or the fact that human emotion and logic still play a massive role in our markets. A sell-off of millions of shares in the airline market could potentially scare humans into selling off their shares in the hotel industry, which in turn could snowball people into selling off their shares in other travel-related companies, which could then affect logistics companies, food supply companies, etc. Take the "Flash Crash" of May 2010 as an example. Towards the end of the trading day, the Dow Jones plunged 1,000 points (more than \$1 trillion in value) before rebounding towards normal levels just 36 minutes later. What caused this crash? A London-based trader named Navinder Singh Sarao first caused the crash and then it became exacerbated by HFT computers. Apparently Sarao used a "spoofing" algorithm that placed an order for thousands of stock index futures contracts betting that the market would fall. Instead of going through with the bet, Sarao was going to cancel the order at the last second and buy the lower priced stocks that were being sold off due to his original bet. Other humans and HFT computers saw this \$200 million bet and took it as a sign that the market was going to tank. In turn, HFT computers began one of the biggest stock sell-offs in history, causing a brief loss of more than \$1 trillion globally. Financial HFT algorithms aren't always correct, either. We view computers as the end-all-be-all when it comes to being correct, but AI is still really just as smart as the humans who programmed it. In 2012, Knight Capital Group experienced a glitch that put them on the verge of bankruptcy. Knight's computers mistakenly streamed thousands of orders per second into the NYSE market causing mass chaos for the company. The HFT algorithms executed an astounding 4 million trades of 397 million shares in only 45 minutes. The volatility created by this computer error led to Knight losing \$460 million overnight and having to be acquired by another firm. Errant algorithms obviously have massive implications for shareholders and the markets themselves, and nobody learned this lesson harder than Knight.



AI IS DETRIMENTAL TO HUMAN RIGHTS AND PRIVACY. THROUGH THE PRIVATE SECTOR AND GOVERNMENTS AROUND THE GLOBE, WE ARE AT AN UNPRECEDENTED LEVEL OF SURVEILLANCE

Dziedzic, 21 -- Dziedzic, M. (2021, September 15). *Urgent action needed over artificial intelligence risks to human rights* | UN news. United Nations. Retrieved June 14, 2022, from <https://news.un.org/en/story/2021/09/1099972>

Urgent action is needed as it can take time to assess and address the serious risks this technology poses to human rights, warned the High Commissioner: "The higher the risk for human rights, the stricter the legal requirements for the use of AI technology should be". Ms. Bachelet also called for AI applications that cannot be used in compliance with international human rights law, to be banned. "Artificial intelligence can be a force for good, helping societies overcome some of the great challenges of our times. But AI technologies can have negative, even catastrophic, effects if they are used without sufficient regard to how they affect people's human rights". Pegasus spyware revelations On Tuesday, the UN rights chief expressed concern about the "unprecedented level of surveillance across the globe by state and private actors", which she insisted was **"incompatible"** with human rights. She was speaking at a Council of Europe hearing on the implications stemming from July's controversy over Pegasus spyware. The Pegasus revelations were no surprise to many people, Ms. Bachelet told the Council of Europe's Committee on Legal Affairs and Human Rights, in reference to the widespread use of spyware commercialized by the NSO group, which affected thousands of people in 45 countries across four continents. 'High price', without action The High Commissioner's call came as her office, OHCHR, published a report that analyses how AI affects people's right to privacy and other rights, including the rights to health, education, freedom of movement, freedom of peaceful assembly and association, and freedom of expression. The document includes an assessment of profiling, automated decision-making and other machine-learning technologies. The situation is "dire" said Tim Engelhardt, Human Rights Officer, Rule of Law and Democracy Section, who was speaking at the launch of the report in Geneva on Wednesday. The situation has "not improved over the years but has become worse" he said. Whilst welcoming "the European Union's agreement to strengthen the rules on control" and "the growth of international voluntary commitments and accountability mechanisms", he warned that "we don't think we will have a solution in the coming year, but the first steps need to be taken now or many people in the world will pay a high price". OHCHR Director of Thematic Engagement, Peggy Hicks, added to Mr Engelhardt's warning, stating "it's not about the risks in future, but the reality today. Without far-reaching shifts, the harms will multiply with scale and speed and we won't know the extent of the problem."



Hill, 20 -- Hill, K. (2020, June 24). *Wrongfully accused by an algorithm*. The New York Times. Retrieved June 8, 2022, from https://www.nytimes.com/2020/06/24/technology/facial-recognition-arrest.html?unlocked_article_code=AAAAAAAAAAAAAAAAACEIPuomT1JKd6J17Vw1cRCfTTMQmqxCdw_Plxftm3iWma3DLDMweiPgYClIG_EPKarska dl43j2fAcVMO7ggQv1uz-hZekV3UQSzvt2EhJEBaW0TmL6EY1kXjdjLTkxqtnjjdHW4l-Nyg-zh4k6MajXvRKLb1Hc-lA4z8Y81Jgr9xXQMwK6RFblK2lQvj-wzRcwwHUd2byKJufLuDBh6KY_GOkmasl9qLrkfDTLDntec6KYCdxFQ Dz_FS3B-5GU_6rBMKY9dffa_f1N7Jp2l0fhGAXdoLYypG5Q0W4HR8r1uurTLohSNo9GkxWaY0X8SmiSqaBvvpzLwEw&smid=url-share

[...] A nationwide debate is raging about racism in law enforcement. Across the country, millions are protesting not just the actions of individual officers, but bias in the systems used to surveil communities and identify people for prosecution. Facial recognition systems have been used by police forces for more than two decades. Recent studies by M.I.T. and the National Institute of Standards and Technology, or NIST, have found that while the technology works relatively well on white men, the results are less accurate for other demographics, in part because of a lack of diversity in the images used to develop the underlying databases. Last year, during a public hearing about the use of facial recognition in Detroit, an assistant police chief was among those who raised concerns. “On the question of false positives — that is absolutely factual, and it’s well-documented,” James White said. “So that concerns me as an African-American male.” This month, Amazon, Microsoft and IBM announced they would stop or pause their facial recognition offerings for law enforcement. The gestures were largely symbolic, given that the companies are not big players in the industry. The technology police departments use is supplied by companies that aren’t household names, such as Vigilant Solutions, Cognitec, NEC, Rank One Computing and Clearview AI. Clare Garvie, a lawyer at Georgetown University’s Center on Privacy and Technology, has written about problems with the government’s use of facial recognition. She argues that low-quality search images — such as a still image from a grainy surveillance video — should be banned, and that the systems currently in use should be tested rigorously for accuracy and bias. “There are mediocre algorithms and there are good ones, and law enforcement should only buy the good ones,” Ms. Garvie said. About Mr. Williams’s experience in Michigan, she added: “I strongly suspect this is not the first case to misidentify someone to arrest them for a crime they didn’t commit. This is just the first time we know about it.” In October 2018, someone shoplifted five watches, worth \$3,800, from a Shinola store in Detroit. Credit... Sylvia Jarrus for The New York Times Mr. Williams’s case combines flawed technology with poor police work, illustrating how facial recognition can go awry. The Shinola shoplifting occurred in October 2018. Katherine Johnston, an investigator at Mackinac Partners, a loss prevention firm, reviewed the store’s surveillance video and sent a copy to the Detroit police, according to their report. Five months later, in March 2019, Jennifer Coulson, a digital image examiner for the Michigan State Police, uploaded a “probe image” — a still from the video, showing the man in the Cardinals cap — to the state’s facial recognition database. The system would have mapped the man’s face and searched for similar ones in a collection of 49 million photos. The state’s technology is supplied for \$5.5 million by a company called DataWorks Plus. Founded in South Carolina in 2000, the company first offered mug shot management software, said Todd Pastorini, a general manager. In 2005, the firm began to expand the product, adding face recognition tools developed by outside vendors. When one of these subcontractors develops an algorithm for recognizing faces, DataWorks attempts to judge its effectiveness by running searches using low-quality images of individuals it knows are present in a system. “We’ve tested a lot of garbage out there,” Mr. Pastorini said. These checks, he added, are not “scientific” — DataWorks does not formally measure the systems’ accuracy or bias. “We’ve become a pseudo-expert in the technology,” Mr. Pastorini said. In Michigan, the DataWorks software used by the state police incorporates components developed by the Japanese tech giant NEC and by Rank One Computing, based in Colorado, according to Mr. Pastorini and a state police spokeswoman. In 2019, algorithms from both companies were included in a federal study of over 100 facial recognition systems that found they were biased, falsely identifying African-American and Asian faces 10 times to 100 times more than Caucasian faces.



FACIAL RECOGNITION TECHNOLOGY SHOULD NOT BE USED AS A “SMOKING GUN.” INACCURACIES WILL PROVE DETRIMENTAL TO MARGINALIZED GROUPS AND THEIR FAMILIES

Hill, 20 -- Hill, K. (2020, June 24). *Wrongfully accused by an algorithm*. The New York Times. Retrieved June 8, 2022, from https://www.nytimes.com/2020/06/24/technology/facial-recognition-arrest.html?unlocked_article_code=AAAAAAAAAAAAAAAAACEIPuomT1JKd6J17Vw1cRCfTTMQmqxCdw_Plxftm3iWma3DLdmweiPgYClIG_EPKarska dl43j2fAcVMO7ggQv1uz-hZekV3UQSzvt2EhJEBaW0TmL6EY1kXjdjLTKxqtnjjdHW4l-Nyg-zh4k6MajXvRKLb1Hc-lA4z8Y81Jgr9xXQMwK6RFblK2lQvj-wzRcwvHUd2byKJufLuDBh6KY_GOkmasl9qLrkfDTLDntec6KYCdxFQ Dz_FS3B-5GU_6rBMKY9dffa_f1N7Jp2l0fhGAXdoLYypG5Q0W4HR8r1uurTLohSNo9GkxWaY0X8SmiSqaBvvpzLwEw&smid=url-share

Rank One’s chief executive, Brendan Klare, said the company had developed a new algorithm for NIST to review that “tightens the differences in accuracy between different demographic cohorts.” After Ms. Coulson, of the state police, ran her search of the probe image, the system would have provided a row of results generated by NEC and a row from Rank One, along with confidence scores. Mr. Williams’s driver’s license photo was among the matches. Ms. Coulson sent it to the Detroit police as an “Investigative Lead Report.” “This document is not a positive identification,” the file says in bold capital letters at the top. “It is an investigative lead only and is not probable cause for arrest.” This is what technology providers and law enforcement always emphasize when defending facial recognition: It is only supposed to be a clue in the case, not a smoking gun. Before arresting Mr. Williams, investigators might have sought other evidence that he committed the theft, such as eyewitness testimony, location data from his phone or proof that he owned the clothing that the suspect was wearing. In this case, however, according to the Detroit police report, investigators simply included Mr. Williams’s picture in a “6-pack photo lineup” they created and showed to Ms. Johnston, Shinola’s loss-prevention contractor, and she identified him. (Ms. Johnston declined to comment.) ‘I guess the computer got it wrong’ The Detroit Detention Center. Mr. Williams was held for 30 hours. Credit... Sylvia Jarrus for The New York Times Mr. Pastorini was taken aback when the process was described to him. “It sounds thin all the way around,” he said. Mr. Klare, of Rank One, found fault with Ms. Johnston’s role in the process. “I am not sure if this qualifies them as an eyewitness, or gives their experience any more weight than other persons who may have viewed that same video after the fact,” he said. John Wise, a spokesman for NEC, said: **“A match using facial recognition alone is not a means for positive identification.”** The Friday that Mr. Williams sat in a Detroit police interrogation room was the day before his 42nd birthday. That morning, his wife emailed his boss to say he would miss work because of a family emergency; it broke his four-year record of perfect attendance. In Mr. Williams’s recollection, after he held the surveillance video still next to his face, the two detectives leaned back in their chairs and looked at one another. One detective, seeming chagrined, said to his partner: “I guess the computer got it wrong.” They turned over a third piece of paper, which was another photo of the man from the Shinola store next to Mr. Williams’s driver’s license. Mr. Williams again pointed out that they were not the same person. Mr. Williams asked if he was free to go. “Unfortunately not,” one detective said. Mr. Williams was kept in custody until that evening, 30 hours after being arrested, and released on a \$1,000 personal bond. He waited outside in the rain for 30 minutes until his wife could pick him up. When he got home at 10 p.m., his five-year-old daughter was still awake. She said she was waiting for him because he had said, while being arrested, that he’d be right back. She has since taken to playing “cops and robbers” and accuses her father of stealing things, insisting on “locking him up” in the living room. Getting help Mr. Williams with his wife, Melissa, and their daughters at home in Farmington Hills, Mich. Credit... Sylvia Jarrus for The New York Times The Williams family contacted defense attorneys, most of whom, they said, assumed Mr. Williams was guilty of the crime and quoted prices of around \$7,000 to represent him. Ms. Williams, a real estate marketing director and food blogger, also tweeted at the American Civil Liberties Union of Michigan, which took an immediate interest. “We’ve been active in trying to sound the alarm bells around facial recognition, both as a threat to privacy when it works and a racist threat to everyone when it doesn’t,” said Phil Mayor, an attorney at the organization. “We know these stories are out there, but they’re hard to hear about because people don’t usually realize they’ve been the victim of a bad facial recognition search.” [...]



Dziedzic, 21 -- Dziedzic, M. (2021, September 15). *Urgent action needed over artificial intelligence risks to human rights* | | UN news. United Nations. Retrieved June 14, 2022, from <https://news.un.org/en/story/2021/09/1099972>

Failure of due diligence According to the report, States and businesses often rushed to incorporate AI applications, failing to carry out due diligence. It states that there have been numerous cases of people being treated unjustly due to AI misuse, such as being denied social security benefits because of faulty AI tools or arrested because of flawed facial recognition software. Discriminatory data The document details how AI systems rely on large data sets, with information about individuals collected, shared, merged and analysed in multiple and often opaque ways. The data used to inform and guide AI systems can be faulty, discriminatory, out of date or irrelevant, it argues, adding that long-term storage of data also poses particular risks, as data could in the future be exploited in as yet unknown ways. “Given the rapid and continuous growth of AI, filling the immense accountability gap in how data is collected, stored, shared and used is one of the most urgent human rights questions we face,” Ms. Bachelet said. The report also stated that serious questions should be raised about the inferences, predictions and monitoring by AI tools, including seeking insights into patterns of human behaviour. It found that the biased datasets relied on by AI systems can lead to discriminatory decisions, which are acute risks for already marginalized groups. “This is why there needs to be systematic assessment and monitoring of the effects of AI systems to identify and mitigate human rights risks,” she added. Biometric technologies An increasingly go-to solution for States, international organizations and technology companies are biometric technologies, which the report states are an area “where more human rights guidance is urgently needed”. These technologies, which include facial recognition, are increasingly used to identify people in real-time and from a distance, potentially allowing unlimited tracking of individuals. The report reiterates calls for a moratorium on their use in public spaces, at least until authorities can demonstrate that there are no significant issues with accuracy or discriminatory impacts and that these AI systems comply with robust privacy and data protection standards. Greater transparency needed The document also highlights a need for much greater transparency by companies and States in how they are developing and using AI. “The complexity of the data environment, algorithms and models underlying the development and operation of AI systems, as well as intentional secrecy of government and private actors are factors undermining meaningful ways for the public to understand the effects of AI systems on human rights and society,” the report says. Guardrails essential “We cannot afford to continue playing catch-up regarding AI – allowing its use with limited or no boundaries or oversight and dealing with the almost inevitable human rights consequences after the fact. “The power of AI to serve people is undeniable, but so is AI’s ability to feed human rights violations at an enormous scale with virtually no visibility. Action is needed now to put human rights guardrails on the use of AI, for the good of all of us,” Ms. Bachelet stressed.



Dolan, 6/8 -- Dolan, C. (2022, June 8). NATO's 2022 strategic concept must enhance digital access and capacities. Just Security. Retrieved June 12, 2022, from <https://www.justsecurity.org/81839/natos-2022-strategic-concept-must-enhance-digital-access-and-capacities/>

This month in Madrid, the North Atlantic Treaty Organization (NATO) will update its Strategic Concept, the principal document that guides the alliance's political-military strategy and collective defense operations. The war in Ukraine has put resilience in the face of Russian aggression front and center, especially in the cyber and information operation domains. Over the years, NATO has digitized and enhanced its security platforms, emphasizing interoperability of systems among its now 30 current member states. If NATO is to become more resilient against advanced persistent threats, hackers, and the maligned states that sponsor them, then the 2022 Strategic Concept must infuse multinational warfighting and deterrence against hybrid threats with methods that facilitate access to data and information sharing on its platforms and across multiple domains, namely in air, cyber, information, land, maritime, and space operations. The 2022 Strategic Concept The Strategic Concept is among NATO's most important documents as it informs alliance planning, resource allocation, and programming based on changes in the threat environment. But the document has not been updated since 2010. The 2010 Strategic Concept, entitled "Active engagement, Modern Defense," contained just one brief sentence about cyber attacks and did not even mention China. It also stated that "Today, the Euro-Atlantic area is at peace," even though Russia had invaded Georgia two years before and the threat of a return to great power competition loomed. To argue that a lot has happened between 2010 and 2022 would be an understatement. Russia's annexation of Crimea and intervention in the Donbas in 2014 and the invasion of Ukraine in 2022 shattered any illusions of a lasting peace with Russia. China's territorial ambitions, economic assertiveness, threats against Taiwan, and military modernization threaten the rules-based order. Emerging technologies – in the form **of hypersonic weapons, artificial intelligence, quantum computing, and machine learning** – have intensified great power competition. The 2022 Strategic Concept should highlight the essential role of technology in collective defense. To build greater digital capacity while also emphasizing resilience, NATO must adopt a new technological orientation on the military strategic level of command, especially within the Allied Command Transformation (ACT) in Norfolk, Virginia and the Allied Command Operations (ACO) in Mons, Belgium. ACT leverages advanced technologies for security and defense in capabilities, procedures, public-private partnerships, civil-military relations, and at NATO's Centers of Excellence. Led by the Supreme Allied Commander Europe, ACO is responsible for collective defense through direction, requirements, planning, and execution at the strategic level. However, the Strategic Concept 2022 should focus less on the emergence of new technologies and more on how NATO's military and civilian personnel use them. ACO and ACT must emphasize greater accessibility to information and data for its multinational warfighters, cyber operators, and civilian professionals. NATO must reach out to experts in the private sector, academia, and non-governmental organizations to harness ways to expand access and emphasize flexibility in multi-domain operations. NATO can do this by providing more grants to private sector partners and establish a new center of excellence on data and information sharing. ACO and ACT should also enable personnel and partners to readily access data and information in DIMEL domains: diplomatic, information/cyber, military, economic, and legal. This would expand the range of measures needed by ACT and ACO to connect and correlate deterrence with evolving hybrid threats.



Dolan, 6/8 -- Dolan, C. (2022, June 8). NATO's 2022 strategic concept must enhance digital access and capacities. Just Security. Retrieved June 12, 2022, from <https://www.justsecurity.org/81839/natos-2022-strategic-concept-must-enhance-digital-access-and-capacities/>

Battlefields and Battlespaces The 2022 Strategic Concept must value partnerships and collaboration with the private sector to develop solutions to pressing collective defense challenges. Technology companies, academics, and start-ups can work with NATO military and civilian personnel to push solutions to multinational warfighters and cyber operators entrusted with protecting alliance members from both conventional military threats and hybrid attacks across multiple domains. This should not be difficult since the world's most advanced capabilities, cutting-edge technology firms, and experts are in NATO member states. NATO still will need advanced stealth F-35s, aircraft carriers, and tanks to wage war on the battlefield. However, NATO also needs digital tools that enable its fighters to readily access information and data in contemporary battlespaces. Whereas "battlefields" highlight land-based operations over others, "battlespaces" are not concerned with a specific arena. Access to data and information is a requirement for multi-domain situational awareness. While it is difficult to predict how future geopolitical events will play out, NATO almost certainly will continue to have to manage threats from an aggressive and revisionist Russia, as well as a rising China. NATO should leverage the Strategic Concept in ways that emphasize access and flexibility against hybrid threats across multiple domains – this is contemporary resilience. Deterring hybrid attacks, whether in the form of cyber intrusions, disinformation through social media platforms, or malignant influence operations, demands that NATO protect its digital infrastructure and allow for greater accessibility to information and data necessary to operate in the gray zone of multi-domain spaces. A "whole-of-alliance" approach grounded on multinational collaboration and coordination, and centered on digital accessibility, will strengthen NATO resilience.

**SOLVENCY EXT.**

USFG REGULATION HAS FAILED AS TECH ADVANCES. TECH AND PRIVATE COMPANIES ARE NEEDED AS THEY TRANSITION TO A GOAL OF SUSTAINABILITY

Thomas, 21 -- Thomas, M. (2021, July 7). *7 dangerous risks of Artificial Intelligence*. Built In. Retrieved June 8, 2022, from <https://builtin.com/artificial-intelligence/risks-of-artificial-intelligence>

Messina is somewhat idealistic about what should happen to help avoid AI chaos, though he's skeptical that it will actually come to pass. Government regulation, he said, isn't a given — especially in light of failures on that front in the social media sphere, whose technological complexities pale in comparison to those of AI. It will take a "very strong effort" on the part of major tech companies to slow progress in the name of greater sustainability and fewer unintended consequences — especially massively damaging ones. "At the moment," he said, "I don't think the onus is there for that to happen." As Messina sees things, it's going to take some sort of "catalyst" to arrive at that point. More specifically, a catastrophic catalyst like war or economic collapse. Though whether such an event will prove "big enough to actually effect meaningful long-term change is probably open for debate." For his part, Ford remains a long-run optimist despite being "very un-bullish" on AI. "I think we can talk about all these risks, and they're very real, but AI is also going to be the most important tool in our toolbox for solving the biggest challenges we face," including climate change. When it comes to the near term, however, his doubts are more pronounced. "We really need to be smarter," he said. "Over the next decade or two, I do worry about these challenges and our ability to adapt to them."



Klosowski, 20 -- Klosowski, T. (2020, July 15). *Facial recognition is everywhere. here's what we can do about it.* The New York Times. Retrieved June 8, 2022, from <https://www.nytimes.com/wirecutter/blog/how-facial-recognition-works/>

The easiest way to see how such regulation might work in practice on a federal level is to look at Illinois's BIPA, which requires consent before an entity can collect and use biometric data (including faceprints) and imposes requirements upon the storage of that data. Consent can be tricky, though. It's one thing for a store to ask if you want to skip showing your ID to enter and another when the store uses this technology to track shoplifters across all its franchise locations. As an example, the EFF's Jennifer Lynch points to a recent case of a business district in London where a company placed cameras in a privately run area that people who worked nearby passed through: "You could see that the business district might say, 'Oh, well, we put up signs,'" Lynch says. "And so people know that when they walk in this area or their face is being recorded and captured, but I don't really believe that people can actually meaningfully consent in that situation. If you are working in that area, you may not have a choice of working somewhere else." When it comes to the government's use of facial recognition, suggested policy approaches diverge. Leong says that although Future of Privacy Forum's main focus is on commercial use of facial recognition, the group would want to see regulation of government use, too. "We would very much like to see overt, intentional regulatory guidance around how the government can and should use facial recognition," she says, "even if it's just things like being really clear about what levels of warrant or probable cause are required for agencies to access it." Other groups, including the EFF, don't think regulation of law enforcement can go far enough. Banning Lynch, along with the EFF, argues that regulation isn't sufficient. "We are pushing for a ban or at least a moratorium at the federal, state, and local level on government use of face recognition," Lynch says. "It is a really game-changing technology and I think we're at a key point in history where we could prevent broad government use of face recognition." Even as facial recognition addresses its diversity problem, there are still too many potential issues concerning how it's used. "The security and policing industries are predicated on this idea that Black people are dangerous," Mutale Nkonde says. "And so when thinking about tools for policing or tools for security, there is going to be this disproportionate deployment against Black people." That is why Nkonde supports banning the software's use outright: "I would want to see a ban around human subjects, just because I think the privacy trade-offs are too huge." Privacy tips for using everyday things with facial recognition Although policy changes, whether in the form of regulation or bans, offer the clearest way forward on a national scale, enacting such changes takes time. Meanwhile, there are smaller but not insignificant ways people interact with facial recognition on a daily basis that are worth thinking deeply about. "I think that the concerning thing and the place where the distinctions sort of blur is that the more we use face recognition, the less we start to think of it, the less we think of it as risky out in the world, we become accustomed to it," says Lynch. "I think it's a slippery slope from using face recognition on your phone to the government using face recognition to track us wherever we go." What about facial recognition in Google Photos or Apple Photos? Photo organization was the first time many people saw facial recognition in action. Apple has made a big show of describing how its facial recognition data in Photos runs on the device (PDF). This technology is more private than a cloud server, but it is also less accurate than cloud-based software. Face grouping in Google Photos can be very accurate, but Google's wide array of services and devices means the company tends to share data liberally across the services it provides.



LARGE COMPANIES WHICH USE FACIAL RECOGNITION TECHNOLOGY POWERED BY AI HAVE STOPPED SELLING THEIR SOFTWARE TO LAW ENFORCEMENTS IN THE US

Klosowski, 20 — Klosowski, T. (2020, July 15). *Facial recognition is everywhere. here's what we can do about it.* The New York Times. Retrieved June 8, 2022, from <https://www.nytimes.com/wirecutter/blog/how-facial-recognition-works/>

Facial recognition—the software that maps, analyzes, and then confirms the identity of a face in a photograph or video—is one of the most powerful surveillance tools ever made. While many people interact with facial recognition merely as a way to unlock their phones or sort their photos, how companies and governments use it will have a far greater impact on people's lives. When it's a device you own or software you use, you may be able to opt out of or turn off facial recognition, but the ubiquity of cameras makes the technology increasingly difficult to avoid in public. Concerns about that ubiquity, amplified by evidence of racial profiling and protester identification, have caused major companies, including Amazon, IBM, and Microsoft, to put a moratorium on selling their software to law enforcement. But as moratoriums expire and the technology behind facial recognition gets better and cheaper, society will need to answer big questions about how facial recognition should be regulated, as well as small questions about which services we're each willing to use and which privacy sacrifices we're each willing to make. How facial recognition software works Most people have seen facial recognition used in movies for decades (video), but it's rarely depicted correctly. Every facial recognition system works differently—often built on proprietary algorithms—but you can sort out the process into three basic types of technology: Detection is the process of finding a face in an image. If you've ever used a camera that detects a face and draws a box around it to auto-focus, you've seen this technology in action. On its own, it isn't nefarious—face detection only focuses on finding a face, not the identity behind it.



Klosowski, 20 -- Klosowski, T. (2020, July 15). *Facial recognition is everywhere. here's what we can do about it.* The New York Times. Retrieved June 8, 2022, from <https://www.nytimes.com/wirecutter/blog/how-facial-recognition-works/>

Analysis (aka attribution) is the step that maps faces—often by measuring the distance between the eyes, the shape of the chin, the distance between the nose and mouth—and then converts that into a string of numbers or points, often called a “faceprint.” Goofy Instagram or Snapchat filters use similar technology (video). Although analysis can suffer from glitches, particularly involving misidentification, that’s generally problematic only when the faceprint is added to a recognition database. Recognition is the attempt to confirm the identity of a person in a photo. This process is used for verification, such as in a security feature on a newer smartphone, or for identification, which attempts to answer the question “Who is in this picture?” And this is where the technology steps into the creepier side of things. The detection phase of facial recognition starts with an algorithm that learns what a face is. Usually the creator of the algorithm does this by “training” it with photos of faces. If you cram in enough pictures to train the algorithm, over time it learns the difference between, say, a wall outlet and a face. Add another algorithm for analysis, and yet another for recognition, and you’ve got a recognition system. The diversity of photos fed into the system has a profound effect on its accuracy during the analysis and recognition steps. For example, if the sample sets mostly include white men—as was the case in the training of early facial recognition systems—the programs will struggle to accurately identify BIPOC faces and women. The best facial recognition software has started to correct for this in recent years, but white males are still falsely matched less frequently (PDF) than other groups; some software misidentifies some Black and Asian people 100 times more often than white men. Mutale Nkonde, fellow of the Digital Civil Society Lab at Stanford and member of the TikTok Content Advisory Council, notes that even if the systems are operating perfectly, issues with gender identification remain: “Labels are typically binary: male, female. There is no way for that type of system to look at non-binary or even somebody who has transitioned.” Once a company trains its software to detect and recognize faces, the software can then find and compare them with other faces in a database. This is the identification step, where the software accesses a database of photos and cross-references to attempt to identify a person based on photos from a variety of sources, from mug shots to photos scraped off social networks. It then displays the results, usually ranking them by accuracy. These systems sound complicated, but with some technical skill, you can build a facial recognition system yourself with off-the-shelf software. A brief history of facial recognition The roots of facial recognition formed in the 1960s, when Woodrow Wilson Bledsoe developed a system of measurements to classify photos of faces. A new, unknown face could then be compared against the data points of previously entered photos. The system wasn’t fast by modern standards, but it proved that the idea had merit. By 1967, interest from law enforcement was already creeping in, and such organizations appear to have funded Bledsoe’s continued research—which was never published—into a matching program. Throughout the ’70s, ’80s, and ’90s, new approaches with catchy names like the “Eigenface approach” (PDF) and “Fisherfaces” improved the technology’s ability to locate a face and then identify features, paving the way for modern automated systems. Facial recognition’s first dramatic shift to the public stage in the US also brought on its first big controversy.



Klosowski, 20 -- Klosowski, T. (2020, July 15). *Facial recognition is everywhere. here's what we can do about it.* The New York Times. Retrieved June 8, 2022, from <https://www.nytimes.com/wirecutter/blog/how-facial-recognition-works/>

In 2001, law enforcement officials used facial recognition on crowds at Super Bowl XXXV. Critics called it a violation of Fourth Amendment rights against unreasonable search and seizure. That year also saw the first widespread police use of the technology with a database operated by the Pinellas County Sheriff's Office, now one of the largest local databases in the country. Skip ahead a few years to 2008, when Illinois's Biometric Information Privacy Act went into effect, becoming the first law of its kind in the US to regulate the unlawful collection and storage of biometric information, including photos of faces. Jennifer Lynch, surveillance litigation director at the Electronic Frontier Foundation, describes BIPA as the model for commercial regulation. "Illinois requires notice and written opt-in consent for the collection of any kind of biometric," she says. "At this point, Illinois is the only state that requires that." The 2010s kickstarted the modern era of facial recognition, as computers were finally powerful enough to train the neural networks required to make facial recognition a standard feature. In 2011, facial recognition served to confirm the identity of Osama bin Laden. In 2014, Facebook publicly revealed its DeepFace photo-tagging software, the same year facial recognition played a key part in convicting a thief in Chicago and the same year Edward Snowden released documents showing the extent to which the US government was collecting images to build a database. In 2015, Baltimore police used facial recognition to identify participants in protests that arose after Freddie Gray was killed by a spinal injury suffered in a police van. Facial recognition first trickled into personal devices as a security feature with Windows Hello and Android's Trusted Face in 2015, and then with the introduction of the iPhone X and Face ID in 2017. Things have ramped up since then: In 2017, President Donald Trump issued an executive order expediting facial recognition usage at US borders (and private airlines have since made their own efforts to incorporate the technology). In 2018, Taylor Swift's security team used facial recognition to identify stalkers, and China rapidly increased its usage. Facial recognition came to Madison Square Garden as a general security measure, and retailers in the US experimented with the tech to track both legitimate shoppers and shoplifters. In 2019, a landlord in New York tried installing it to replace keys, and several schools attempted the same. Today, a handful of cities—San Francisco, Oakland, and Berkeley in California, plus Boston and Somerville in Massachusetts—have banned facial recognition usage by government entities. The country has also seen the first known case of a false positive leading to an arrest in the US. After Black Lives Matter police-brutality protests started in June, several large facial recognition vendors, including Amazon, IBM, and Microsoft, put a halt on selling their technology to law enforcement. However, other, new players have entered the arena. Clearview AI made news in early 2020 when The New York Times revealed that the company regularly ran its recognition software against a database of photos scraped from sources across the internet, including social media, news sites, and employment sites—which Wirecutter, and many others, were able to confirm with testing—in a process that it used to identify suspects. In May 2020, the ACLU announced a lawsuit against Clearview AI in Illinois state court alleging that it violated the privacy rights of Illinois residents under BIPA. Clearview AI is an outlier only in that it has faced public scrutiny: Equally less ethical software companies exist—companies that will sell their software to local law enforcement, usually with no oversight or public scrutiny into where the photos come from or how the identification algorithms work.



ERROR VIA THE CRIMINAL JUSTICE SYSTEM SPURS PRIVACY CONCERNS. FALSE ARRESTS OR LACK THEREOF BREWS MISTRUST IN THE TECHNOLOGY

Klosowski, 20 — Klosowski, T. (2020, July 15). *Facial recognition is everywhere. here's what we can do about it.* The New York Times. Retrieved June 8, 2022, from <https://www.nytimes.com/wirecutter/blog/how-facial-recognition-works/>

The arguments for and against facial recognition Proponents of facial recognition suggest that the software is useful because alongside identifying suspects, it can monitor known criminals and help identify child victims of abuse. In crowds, it could monitor for suspects at large events and increase security at airports or border crossings. The most long-running type of facial recognition software runs a photo through a government-controlled database, such as the FBI's database of over 400 million photos, which includes driver's licenses from some states, to identify a suspect. Local police departments use a variety of facial recognition software, often purchased from private companies. There's a long list of benefits facial recognition can offer outside of law enforcement, adding convenience or security to everyday things and experiences. Facial recognition is helpful for organizing photos, useful in securing devices like laptops and phones, and beneficial in assisting blind and low-vision communities. It can be a more secure option for entry into places of business, fraud protection at ATMs, event registration, or logging in to online accounts. Advertising and commercial applications of facial recognition promise a wide array of supposed benefits, including tracking customer behavior in a store to personalize ads online. Brenda Leong, senior counsel and director of artificial intelligence and ethics at Future of Privacy Forum, suggested in an interview that proponents point to facial recognition as a replacement for loyalty programs or gated access: "You just walk through a set of cameras and all those things happen very seamlessly, sports arenas, event venues, amusement parks, all those places either are using or would have ideas of ways to use it similarly." Opponents don't think these benefits are worth the privacy risks, nor do they trust the systems or the people running them. The first point of contention lies in the act of collection itself—it's very easy for law enforcement to collect photos but nearly impossible for the public to avoid having their images taken. Mug shots, for example, happen upon arrest but before conviction. Error rates in recognition are also problematic, both in a false-positive sense, where an innocent person is falsely identified, and a false-negative sense, where a guilty person isn't identified. The facial recognition software that law enforcement agencies use isn't currently available for public audit, and the algorithms that power the detection and identification software are often closed-box proprietary systems that researchers can't investigate.



Klosowski, 20 -- Klosowski, T. (2020, July 15). *Facial recognition is everywhere. here's what we can do about it.* The New York Times. Retrieved June 8, 2022, from <https://www.nytimes.com/wirecutter/blog/how-facial-recognition-works/>

When the public doesn't know how these facial recognition systems work or how accurate they are, the public doesn't know whether these systems are being used appropriately, especially in law enforcement. Joseph Flores, a software developer who in his free time uses machine learning for art projects (disclosure: I've worked on related artistic projects with Flores, for fun, not for profit), explained to me how he often intentionally biases his data sets to produce the results he wants, something law enforcement could also do: “You could do the same with your law enforcement facial recognition data to make sure that your friends were unrecognizable and your enemies were misidentified as criminals.” Flores adds, “It's hard to challenge the legality or the reliability of math that you can't review. Especially with the data scale we're talking about. With no review everything is falsifiable and just modern phrenology.”

Another growing issue is law enforcement's interest in real-time recognition in live video feeds or police body-cam footage. But even cities that enthusiastically moved forward with the technology, such as Orlando, Florida—where the police department used Amazon's Rekognition software to attempt to identify suspects in real time from video streams—have dialed those efforts back after the technology failed to live up to expectations. But just because real-time facial recognition still suffers from hiccups on a large scale in live testing doesn't mean it won't become widespread in the future. The idea is so appalling to some communities that the practice is already temporarily banned in California, Oregon, and New Hampshire. The future of facial recognition and regulation Generally speaking, the future of facial recognition can take any of three possible forms: no regulation at all, some regulation, and banning. No regulation The Black Mirror episodes illustrating a world devoid of facial recognition regulation write themselves. Brenda Leong provided a few examples: “It's very easy to create very Orwellian futures, where things are tracking you everywhere you go by your face because cameras are everywhere. If you're a student it could be literally watching whether you're focusing on your work versus daydreaming. If you're an employee, monitoring your engagement on your computer or telling whether you wandered off somewhere else.” **The list of surveillance possibilities is nearly endless,** with China's “Social Credit Score” or the London police force's use of facial recognition cameras in real time offering a glimpse of one particularly grim reality. Regulation As of this writing, there's one proposed US law on a federal level banning police and FBI use of facial recognition, as well as another that allows exceptions with a warrant. Still another bill requires businesses to ask consent before using facial recognition software publicly, and yet another bans its use in public housing. Although facial recognition is certainly having a moment, it's still unclear which of these bills, if any, will have enough support to become laws. When anyone talks about regulating facial recognition, they need to divide the idea into two parts: regulating commercial use and regulating government use, including that of law enforcement. For commercial use, Leong stresses, the main thrust of regulation concerning any commercial feature—a loyalty program, theme park VIP access, or whatever else—should be consent. Facial recognition “should never be the default,” she says. “It should never be part of the standard terms of service or privacy policy. And it should never be like the thing that happens that you have to then go opt out of.”



Klosowski, 20 -- Klosowski, T. (2020, July 15). *Facial recognition is everywhere. here's what we can do about it.* The New York Times. Retrieved June 8, 2022, from <https://www.nytimes.com/wirecutter/blog/how-facial-recognition-works/>

In 2016, Google was sued in Illinois for its use of facial recognition, but that suit was later dismissed. In 2020, a new class action suit alleges a similar offense. Although the ability to organize photos by faces using the facial recognition feature in a photos app offers quantifiable benefits, there is a privacy trade-off to consider. It's difficult to know exactly how a company might misuse your data; this was the case with the photo storage company Ever, whose customers trained the Ever AI algorithm without realizing it. You can disable face grouping in Google Photos. You can't turn the corresponding feature off in Apple's Photos app, but if you don't actively go in and link a photo to a name, the recognition data never leaves your device. What about Facebook? Facebook likely has the largest facial data set ever assembled, and if Facebook has proven anything over the years, it's that people shouldn't trust the company to do the right thing with the data it collects. Facebook recently agreed to pay \$550 million to settle a lawsuit in Illinois over its photo tagging system. Here's how to opt out. What about unlocking a phone or computer? As the features work now, face unlock typically happens only on the device itself, and that data is never uploaded to a server or added to a database. What about facial recognition in home security cameras? The systems behind security cameras lack clear consent as they record and opt-in people automatically, often in defiance of local privacy laws, an ethical problem many people neglect to consider. Right now, only a handful of home security cameras include facial recognition, including Wirecutter's smart doorbell upgrade pick, Google's Nest Hello. Face detection on Nest cameras is off by default, however. More worrisome to privacy advocates is the potential inclusion of facial recognition with Ring cameras, a system that shares data with police through its Neighbors app. Do you need to worry about those goofy face apps that pop up once a year or so? The most recent app to break through in this arena was FaceApp, which gained popularity by allowing people to age themselves. Although the company says it doesn't use the app to train facial recognition software, it's difficult to know what might happen with the data the app collects if the company gets sold. The same goes for whatever the next version of FaceApp is. It's best to be wary of this type of software. Can facial recognition identify you if you're wearing a mask? It's not likely right now but may be in the future. One company in China was able to get facial recognition working on 95% of mask wearers, but this specific software was designed for small-scale databases of around 50,000 employees. Companies are scrambling to solve this problem. Where society goes from here promises to be a mixture of policy and tweaks to people's personal habits, but the conversation concerning the technology likely isn't going anywhere for a long time. Like any technology, facial recognition is itself just software, but as Mutale Nkonde notes, how society uses it is what matters: "It's the way the tool impacts our civil and human rights that is my point of intervention, because I think that all technology is agnostic."



CHINA EXT.

THIS INEVITABLY LEADS TO CHINESE MILITARY ADVANCEMENT AND INCREASED COMPETITIVENESS IN THE NAME OF “NATIONAL SECURITY”

Creemers et al, 2017 -- Creemers, R., Triolo, P., Webster, G., & Kania, E. B. (2017, August 1). *China's plan to 'lead' in AI: Purpose, prospects, and Problems*. New America. Retrieved June 13, 2022, from <https://www.newamerica.org/cybersecurity-initiative/blog/chinas-plan-lead-ai-purpose-prospects-and-problems/>

Based on these foundations for its innovation ambitions, China plans to pursue cutting-edge advances in a category of critical next-generation AI technologies in order to “occupy the commanding heights” of AI science and technology. The plan calls for progress in new “AI 2.0” technologies, including big data intelligence, cross-media intelligence, swarm intelligence, hybrid-augmented intelligence (e.g., human-machine symbiosis or brain-computer collaboration), swarm intelligence, and autonomous intelligent systems, among others. China appears to be particularly focused on approaches that could enable paradigmatic changes in AI, such as high-level machine learning (e.g., self-adaptive learning or autonomous learning), brain-inspired AI, and quantum-accelerated machine learning. The trajectory of this ambitious agenda for research and development remains to be seen, but it is clear that China aspires to lead the world in next-generation AI and could devote extensive resources, including massive amounts of funding, data, and human capital, to this endeavor. Through this AI 2.0 agenda, the Chinese government plans to leverage its rise in AI to enhance national competitiveness, while bolstering its capacity to ensure state security and national defense. The Chinese Communist Party will certainly seek to direct the development of AI in accordance with the interests and imperatives of the Party-state. Consequently, AI will have a range of applications in the domain of “social governance,” including to protect public security and social stability. Certain of these AI applications are already being deployed, including the use of facial recognition to track down criminals and dissidents and even attempts to predict future criminal behavior based on human behavior, reminiscent of Minority Report. According to the plan, the Chinese government will leverage AI to create systems for intelligent monitoring and early warning and control of potential (or perceived) threats. Concurrently, the Chinese leadership wants to ensure that advances in AI can be leveraged for national defense, through a national strategy for military-civil fusion (军民融合). According to the plan, resources and advances will be shared and transferred between civilian and military contexts. This will involve the establishment and normalizing of mechanisms for communication and coordination among scientific research institutes, universities, enterprises, and military industry. Utilizing such approaches, China intends to apply new-generation AI as a “powerful support” to command decision-making, military deduction, defense equipment, and other areas. In practice, repurposing breakthroughs in AI across domains will not necessarily be straightforward. Nonetheless, the PLA will seek to leverage the available resource and advances to enable defense innovation and enhance its capabilities in future “intelligentized” (智能化) warfare, potentially changing paradigms of military power.



AI R&D, DATA TRAINING, AND TALENT ACQUISITION WILL BE DOMINATED BY THE PRIVATE SECTOR

Creemers et al, 2017 -- Creemers, R., Triolo, P., Webster, G., & Kania, E. B. (2017, August 1). *China's plan to 'lead' in AI: Purpose, prospects, and Problems*. New America. Retrieved June 13, 2022, from <https://www.newamerica.org/cybersecurity-initiative/blog/chinas-plan-lead-ai-purpose-prospects-and-problems/>

Meanings of Leadership in AI Government scientists, S&T bureaucrats, and planners are unlikely to substantively lead China's AI development when private companies are the ones developing crucial data resources and are much more able to attract and pay for top AI talent. Existing momentum in the private sector will already doubtless make Chinese efforts among the world leaders in several types of AI applications by 2030, though it may not be a government plan that makes the difference. The question remains what that leadership means, as many AI applications are developed based on culturally, linguistically, and geographically bounded data. For instance, if a Chinese company develops natural language processing technology that gives Chinese users a level of capability unmatched globally, that doesn't mean it can necessarily market the same service in other languages. Other more mundane challenges include ownership of high resolution digital maps that will be required for autonomous driving. At very least, the notion of one nation leading in AI generally will be complicated by the field's diversity of technical and social challenges, as well as the different inputs necessary for different applications—not to mention that a comprehensively leading effort will by definition take place across borders. The plan's recognition of the need for regulatory, legal, and ethical principles for AI development and use does, however, represent an uncommonly foresighted approach. Of course, the Chinese government's approach to AI regulation, ethics, and economic adjustment will reflect its broader model of governance and ideology. Thus it will be crucial for other jurisdictions, for instance the United States and the EU, to develop regulatory, ethical, and developmental approaches that reflect their own values. Part III: Changing Paradigms Through Independent Innovation and "Next Generation AI" By Elsa B. Kania The Chinese leadership sees technological innovation, particularly in AI, as a core aspect of international competition. Beyond informatization, China is embarking upon an agenda of "intelligentization" (智能化), seeking to take advantage of the transformative potential of AI throughout society, the economy, government, and the military. Through this new plan, China intends to pursue "indigenous innovation" in the "strategic frontier" technology of AI in furtherance of a national strategy for innovation-driven development. If this were only rhetoric, such a focus on innovation might remain aspirational. However, this plan includes an extensive and detailed agenda, with sustained focus and significant funding, to build up China's capability in innovation capability to enable advances in next-generation AI technologies. Looking forward, China may have the potential to lead the world in this new AI revolution, potentially surpassing the United States in the process. Despite its remarkable rise in AI, this plan candidly acknowledges several shortcomings in China's current capacity relative to the cutting edge of the field. To date, China has not produced major original results in AI and has lagged in critical components, such as high-performance chips for machine learning. There has not been a systematic high-level design for research and development. Chinese research institutions and enterprises have yet to establish influence at the international level. Meanwhile, the aggressive attempts of major Chinese tech companies to recruit leading AI talent abroad is seemingly symptomatic of the gap between the available talent within China and the demand for it. While China remains in the process of building up indigenous capacity, the plan calls for the leveraging of "international innovation resources."

**RUSSIA/UKRAINE EXT.**

RUSSIA ALREADY AGAINST NATO NOW — INCREASED RUSSIAN AGGRESSION LOOMS WAR

Blinken, 1/20 — Blinken, A. J. (2022, January 20). *The stakes of Russian aggression for Ukraine and beyond* - united states department of state. U.S. Department of State. Retrieved June 15, 2022, from <https://www.state.gov/the-stakes-of-russian-aggression-for-ukraine-and-beyond/>

Good afternoon. First, let me say how honored I am by the presence of so many friends, colleagues, leaders across different communities here in Germany, and also leaders in the partnership that links our two countries. I'm grateful to all of you for being here, grateful for this opportunity as well to be at the Academy of Sciences and Humanities. I heard a little bit from Sigmar about the history, briefly walked the hallways, and I very much appreciate this hospitality. But it's an institution with an extraordinary tradition of scholarship, discovery stretching back more than 300 years. And I understand that, among other luminaries, Albert Einstein was a member here, so I should probably let you know that my remarks today will include very little about astrophysics, which will be to everyone's benefit. I want to thank all the institutions that are cohosting us, including Atlantik-Brücke. By the way, my own history with the Brücke, the bridge, goes back well more than 20 years. I remember very well spending time with visiting colleagues from Germany during the Clinton administration. But it's pleasure to be with you, the German Marshall Fund, the Aspen Institute, the American Council on Germany. And I can't not acknowledge a great friend, colleague going back to my university days, the Clinton administration, the Obama administration, Dan Benjamin. It's wonderful to see you as well. Over the years, these organizations have helped build, strengthen, and deepen the ties between our countries. One of the markers of a strong democracy is a robust, independent civil society, and I'm grateful to our cohosts for their contributions to democracy on both sides of the Atlantic and, again, for bringing us together today. So as Sigmar said, and as all of you know, I have come to Berlin at a moment of great urgency for Europe, for the United States, and, I would argue, for the world. Russia is continuing to escalate its threat toward Ukraine. We've seen that again in just the last few days with increasingly bellicose rhetoric, building up its forces on Ukraine's borders, including now in Belarus. Russia has repeatedly turned away from agreements that have kept the peace across the continent for decades. And it continues to take aim at NATO, a defensive, voluntary alliance that protects nearly a billion people across Europe and North America, and at the governing principles of international peace and security that we all have a stake in defending. Those principles, established in the wake of two world wars and a cold war, reject the right of one country to change the borders of another by force; to dictate to another the policies it pursues or the choices it makes, including with whom to associate; or to exert a sphere of influence that would subjugate sovereign neighbors to its will. To allow Russia to violate those principles with impunity would drag us all back to a much more dangerous and unstable time, when this continent and this city were divided in two, separated by no man's lands, patrolled by soldiers, with the threat of all-out war hanging over everyone's heads. It would also send a message to others around the world that these principles are expendable, and that, too, would have catastrophic results.



Blinken, 1/20 -- Blinken, A. J. (2022, January 20). *The stakes of Russian aggression for Ukraine and beyond* - united states department of state. U.S. Department of State. Retrieved June 15, 2022, from <https://www.state.gov/the-stakes-of-russian-aggression-for-ukraine-and-beyond/>

That's why the United States and our allies and partners in Europe have been so focused on what's happening in Ukraine. It's bigger than a conflict between two countries. It's bigger than Russia and NATO. It's a crisis with global consequences, and it requires global attention and action. Here today, among this rapidly unfolding situation, I'd like to try to cut through to the facts of the matter. To begin, Russia claims that this crisis is about its national defense, about military exercises, weapons systems, and security agreements. Now, if that's true, we can resolve things peacefully and diplomatically. There are steps we can take – the United States, Russia, the countries of Europe – to increase transparency, reduce risks, advance arms control, build trust. We've done this successfully in the past and we can do it again. And, indeed, it's what we set out to do last week in the discussions that we put forward at the Strategic Stability Dialogue between the United States and Russia, at the NATO-Russia Council, and at the OSCE. At those meetings and many others, the United States and our European allies and partners have repeatedly reached out to Russia with offers of diplomacy in a spirit of reciprocity. So far, our readiness to engage in good faith has been rebuffed, because in truth this crisis is not primarily about weapons or military bases. It's about the sovereignty and self-determination of Ukraine and all states. And at its core, it's about Russia's rejection of a post-Cold War Europe that is whole, free, and at peace. For all our profound concerns with Russia's aggression, provocations, political interference – including against the United States – the Biden administration has made clear our willingness to pursue a more stable, predictable relationship; to negotiate arms control agreements, like the renewal of New START, and launch our Strategic Stability Dialogue; to pursue common action to address the climate crisis and work in common cause to revive the Iran nuclear deal. And we appreciate how Russia has engaged with us in these efforts. And despite Moscow's reckless threats against Ukraine and dangerous military mobilization – despite its obfuscation and disinformation – the United States, together with our allies and partners, have offered a diplomatic path out of this contrived crisis. That's why I've returned to Europe – Ukraine yesterday, Germany here today, Switzerland tomorrow, where I'll meet with Russian Foreign Minister Lavrov and once again seek diplomatic solutions. The United States would greatly prefer those to be the case, and certainly prefer diplomacy to the alternatives. We know our partners in Europe feel the same way. So do people and families across the continent, because they know that they will bear the greatest burden if Russia rejects diplomacy. And we look to countries beyond Europe, to the international community as a whole to make clear the costs to Russia if it seeks conflict, and to stand up for all the principles that protect all of us. So let's look plainly at what's at stake right now, who will actually be affected, and who is responsible. In 1991, millions of Ukrainians went to the polls to say that Ukraine would no longer be ruled by autocrats but would govern itself. In 2014, the Ukrainian people stood up to defend their choice for a democratic and European future. They've been living under the shadow of Russian occupation in Crimea and aggression in Donbas ever since.



Blinken, 1/20 — Blinken, A. J. (2022, January 20). *The stakes of Russian aggression for Ukraine and beyond* - united states department of state. U.S. Department of State. Retrieved June 15, 2022, from <https://www.state.gov/the-stakes-of-russian-aggression-for-ukraine-and-beyond/>

The war in eastern Ukraine, orchestrated by Russia with proxies that it leads, trains, supplies, and finances – well, that's killed more than 14,000 Ukrainians. Thousands more have been wounded. Entire towns have been destroyed. Nearly one and a half million Ukrainians have fled their homes to escape the violence. For Ukrainians in Crimea and the Donbas, the repression is acute. Russia blocks Ukrainians from crossing the line of contact, cutting them off from the rest of the country. Hundreds of Ukrainians are being held as political prisoners by Russia and its proxies. Hundreds of families don't know if their loved ones are alive or dead. And the humanitarian needs are growing. Nearly 3 million Ukrainians, including a million elderly people and half a million children, urgently need food, shelter, and other life-saving assistance. But of course, even Ukrainians who live far away from the fighting are affected by it. This is their country; these are their fellow citizens. And nowhere in Ukraine are people free from Russia's malign activities. Moscow has sought to undermine Ukraine's democratic institutions, interfered in Ukraine's politics and elections, blocked energy and commerce to intimidate Ukraine's leaders and pressure its citizens, used propaganda and disinformation to sow mistrust, launched cyber attacks on the country's critical infrastructure. The campaign to destabilize Ukraine has been relentless. And now Russia is poised to go even further. The human toll of renewed aggression by Russia would be by many magnitudes higher than what we've seen to date. Russia justifies its actions by claiming that Ukraine somehow poses a threat to its security. This turns reality on its head. Whose troops are surrounding whom? Which country has claimed another's territory through force? Which military is many times the size of the other? Which country has nuclear weapons? Ukraine isn't the aggressor here; Ukraine is just trying to survive. No one should be surprised if Russia instigates a provocation or incident and then tries to use it to justify military intervention, hoping that by the time the world realizes the ruse it'll be too late. There's been a lot of speculation about President Putin's true intentions, but we don't actually have to guess. He's told us repeatedly. He's laying the groundwork for an invasion because he doesn't believe that Ukraine is a sovereign nation. He said it flat out to President Bush in 2008, and I quote, "Ukraine isn't a real country." He said in 2020, and I quote, "Ukrainians and Russians are one and the same people." Just a few days ago, the Russian ministry of foreign affairs tweeted in celebration of the anniversary of Ukraine and Russia's unification in the year 1654. That's a pretty unmistakable message this week of all weeks. And so the stakes for Ukraine come more fully into view. This is not only about a possible invasion and war. It's about whether Ukraine has a right to exist as a sovereign nation. It's about whether Ukraine has a right to be a democracy. This hasn't stopped with Ukraine. All the former Soviet socialist republics became sovereign nations in 1990 and 1991. One of them is Georgia. Russia invaded it in 2008. Thirteen years later, nearly 300,000 Georgians are still displaced from their homes. Another is Moldova. Russia maintains troops and munitions there against the will of its people. If Russia invades and occupies Ukraine, what's next? Certainly, Russia's efforts to turn its neighbors into puppet states, to control their activities, to crack down on any spark of democratic expression will intensify. Once the principles of sovereignty and self-determination are thrown out, you revert to a world in which the rules we shaped together over decades erode and then vanish.



COOPERATIVE EFFORTS BETWEEN NATO AND RUSSIA HAVE FAILED HISTORICALLY — ONLY AMPS RUSSIAN AGGRESSION AND DESTRUCTIVE WEAPON DEVELOPMENT

Blinken, 1/20 — Blinken, A. J. (2022, January 20). *The stakes of Russian aggression for Ukraine and beyond* - united states department of state. U.S. Department of State. Retrieved June 15, 2022, from <https://www.state.gov/the-stakes-of-russian-aggression-for-ukraine-and-beyond/>

And that emboldened some governments to do whatever it takes to get whatever they want, even if that means shutting down another country's internet, cutting off heating oil in the dead of winter, or sending in tanks — all tactics Russia has used against other countries in recent years. That's why governments and citizens everywhere should care about what's happening in Ukraine. It may seem like a distant regional dispute or yet another example of Russian bullying, but at stake, again, are principles that have made the world safer and more stable for decades. Now alternatively, **Russia says the problem is NATO.** On its face, that's absurd. NATO didn't invade Georgia; NATO didn't invade Ukraine. Russia did. NATO is a defensive Alliance with no aggressive intent toward Russia. To the contrary, efforts by NATO to engage Russia have gone on for years, and unfortunately, been rejected. For example, in the NATO Russia Founding Act, which was intended to build trust and increase consultations and cooperation, NATO pledged to significantly reduce its military strength in Eastern Europe. And it's done just that. Russia pledged to exercise similar restraint in its conventional force deployments in Europe. Again, instead, it invaded two countries. Russia says that NATO is encircling Russia. In fact, only 6 percent of Russia's borders touch NATO countries. Compare that to Ukraine, which is now genuinely being encircled by Russian troops. In the Baltic countries and Poland, there are around 5,000 NATO troops who aren't from those countries, and their presence is rotational, not permanent. Russia has put at least 20 times as many on Ukraine's borders. President Putin says that NATO is, and I quote, "parking missiles on the porch of our house." But it's Russia that has developed ground-launched, intermediate-range missiles that can reach Germany and nearly all NATO European territory despite Russia being a party to the INF Treaty that prohibited these missiles. In fact, Russia's violation led to the termination of that treaty, which has left us all less safe. It's also worth noting that though Russia is not a member of NATO, it, like many non-NATO countries, has actually benefited from the peace, stability, and prosperity that NATO has helped make possible. Many of us remember vividly the tensions and fears of the Cold War era. The steps that the Soviet Union and the West took toward each other over those years to build understanding and establish agreed-upon rules for how our countries would act were welcomed by people everywhere because they turned down the heat and made military conflict less likely. Those breakthroughs are the result of a great deal of hard work by people on all sides. Now we're seeing that hard work come undone. For example, in 1975, all OSCE countries, including Russia, signed the Helsinki Final Act, which established 10 guiding principles for international behavior, including respect for national sovereignty, refraining from the threat or use of force, the inviolability of frontiers, the territorial integrity of states, the peaceful settlement of disputes, and non-intervention in internal affairs. Russia has since violated every single one of those principles in Ukraine and has repeatedly made clear its disdain for them.



RUSSIA HAS REPEATEDLY GONE BACK ON PEACE PACTS AND TREATIES. THEY WILL CONTINUE R&D OF AI DESPITE OF NATO REGULATIONS OR OVERSIGHT. THIS LEADS TO INTERNATIONAL CONFLICT

Blinken, 1/20 -- Blinken, A. J. (2022, January 20). *The stakes of Russian aggression for Ukraine and beyond* - united states department of state. U.S. Department of State. Retrieved June 15, 2022, from <https://www.state.gov/the-stakes-of-russian-aggression-for-ukraine-and-beyond/>

In 1990, the OSCE countries, including Russia, agreed to the Vienna Document, a set of confidence- and security-building measures to increase transparency and predictability about military activities, including military exercises. Now, Russia selectively follows those provisions. For example, it holds large-scale military exercises that it claims are exempt from the notification and observation requirements of the Vienna Document because they're conducted without prior notice to the troops involved. Last fall, Russia conducted military exercises in Belarus with more than 100,000 troops. It's impossible that those exercises were no notice. And Moscow has failed to provide information on its military forces in Georgia, to notify the OSCE of its massive troop buildup around Ukraine last spring, to answer Ukraine's questions about what it was doing, all of which are required under that 1990 agreement. In 1994, in a pact known as the Budapest Memorandum, Russia, the United States and Britain committed to, and I quote, "respect the independence and sovereignty and the existing borders of Ukraine and to refrain from the threat or use of force against" the country. Those promises helped persuade Ukraine to give up their nuclear arsenal inherited after the dissolution of the USSR and which was then the third largest nuclear arsenal in the world. Well, we need only ask the people living in Crimea and Donbas what happened to those pledges. There are many more examples I could cite. They all support the same conclusion: One country has repeatedly gone back on its commitments and ignored the very rules it agreed to despite others working hard to bring it along at every step. That country is Russia. Of course, Russia is entitled to protect itself, and the United States and Europe are prepared to discuss Russia's security concerns and how we can address them in a reciprocal way. Russia has concerns about its security and actions that it says the United States and Europe and NATO are taking that somehow threaten that security. We have profound concerns about the actions that Russia is taking that threaten our security. We can talk about all of that. But we will not treat the principles of sovereignty or territorial integrity enshrined in the UN Charter, affirmed by the UN Security Council, as negotiable. And if I could speak to the Russian people, I would say to them you deserve to live with security and dignity like all people everywhere, and no one – not Ukraine, not the United States, not NATO or its members – is seeking to jeopardize that. But what really risks your security is a pointless war with your neighbors in Ukraine with all the costs that come with it, most of all for the young people who will risk or even give their lives to it. At a time when COVID is running throughout the planet, we have a climate crisis, we need to rebuild the global economy, all of which demand so much of our attention and resources, is this really what you need – a violent conflict that will likely drag on? Would that actually make your lives more secure, more prosperous, more full of opportunity? And just think of what a great nation like Russia could achieve if it dedicated its resources, especially the remarkable talent of its human resources, its people, toward the most significant challenges of our time. We in the United States, our partners in Europe, we would welcome that.



Blinken, 1/20 -- Blinken, A. J. (2022, January 20). *The stakes of Russian aggression for Ukraine and beyond* - united states department of state. U.S. Department of State. Retrieved June 15, 2022, from <https://www.state.gov/the-stakes-of-russian-aggression-for-ukraine-and-beyond/>

Tomorrow I'll meet with Foreign Minister Lavrov and I'll urge that Russia find its way back to the agreements it swore to over the decades and to working with the United States and our allies and partners in Europe to write a future that can ensure our mutual security but also make clear that that possibility will be extinguished by Russian aggression against Ukraine, which would also do the very thing Moscow complains about: bolster the NATO defensive alliance. These are difficult issues we're facing. Resolving them won't happen quickly. I certainly don't expect we'll solve them in Geneva tomorrow. But we can advance our mutual understanding. And that, combined with de-escalation of Russia's military buildup on Ukraine's borders – that can turn us away from this crisis in the weeks ahead. At the same time, the United States will continue to work with our allies and partners in NATO, the European Union, the OSCE, the G7, the United Nations, throughout the international community to make clear that there are two paths before Russia: the path to diplomacy that can lead to peace and security; and the path of aggression that will lead only to conflict, severe consequences, international condemnation. The United States and our allies will continue to stand with Ukraine and to stand ready to meet Russia on either path. It's no accident that I'm offering these thoughts here in Berlin. Perhaps no place in the world experienced the divisions of the Cold War more than this city. Here, President Kennedy declared all free people citizens of Berlin. Here, President Reagan urged Mr. Gorbachev to tear down that wall. It seems a time that President Putin wants to return to that era. We hope not. But if he chooses to do so, he'll be met with the same determination, the same unity that past generations of leaders and citizens brought to bear to advance peace, to advance freedom, to advance human dignity across Europe and around the world. Thanks so much for listening.



MORE NEG ANSWERS

IMPACT - NUCLEAR WAR

1) Overinvestment in AI can lead to overdependence, leading to fatally unintentional consequences such as nuclear war.

Kallenborn, 22 (Zachary Kallenborn, Zachary Kallenborn is a research affiliate with the Unconventional Weapons and Technology Division of the National Consortium for the Study of Terrorism and Responses to Terrorism (START), a policy fellow at the Schar School of Policy and Government, a US Army Training and Doctrine Command "Mad Scientist," and national security consultant, 2-1-2022, accessed on 6-13-2022, Bulletin of the Atomic Scientists, "Giving an AI control of nuclear weapons: What could possibly go wrong? - Bulletin of the Atomic Scientists", <https://thebulletin.org/2022/02/giving-an-ai-control-of-nuclear-weapons-what-could-possibly-go-wrong/>)

How autonomous nuclear weapons could go wrong. The huge problem with autonomous nuclear weapons, and really all autonomous weapons, is error. Machine learning-based artificial intelligences—the current AI vogue—rely on large amounts of data to perform a task. Google's AlphaGo program beat the world's greatest human go players, experts at the ancient Chinese game that's even more complex than chess, by playing millions of games against itself to learn the game. For a constrained game like Go, that worked well. But in the real world, data may be biased or incomplete in all sorts of ways. For example, one hiring algorithm concluded being named Jared and playing high school lacrosse was the most reliable indicator of job performance, probably because it picked up on human biases in the data.

In a nuclear weapons context, a government may have little data about adversary military platforms; existing data may be structurally biased, by, for example, relying on satellite imagery; or data may not account for obvious, expected variations such as imagery in taken during foggy, rainy, or overcast weather.

2) AND, WARS UTILIZING AUTONOMOUS WEAPONS ARE MORE LIKELY TO ESCALATE DUE TO PREEMPTIVE ATTACKS.

Laird, 16 (Burgess Laird, Burgess Laird is a Senior International Defense Researcher with the RAND Corporation and an adjunct instructor in the M.A. in Global Security Studies at Johns Hopkins University, 12-8-2016, accessed on 6-9-2022, Rand, "The Risks of Autonomous Weapons Systems for Crisis Stability and Conflict Escalation in Future U.S.-Russia Confrontations", <https://www.rand.org/blog/2020/06/the-risks-of-autonomous-weapons-systems-for-crisis.html>)

First, a state facing an adversary with AWS capable of making decisions at machine speeds is likely to fear the threat of sudden and potent attack, a threat that would compress the amount of time for strategic decisionmaking. The posturing of AWS during a crisis would likely create fears that one's forces could suffer significant, if not decisive, strikes. These fears in turn could translate into pressures to strike first—to preempt—for fear of having to strike second from a greatly weakened position. Similarly, within conflict, the fear of losing at machine speeds would be likely to cause a state to escalate the intensity of the conflict possibly even to the level of nuclear use.