

主催：全模研

第9回 全国中高教育 模擬国連大会 AJEMUN

～ 高校生の高校生による高校生のための模擬国連大会～

『AIと軍事』

AI and the military

8/4(月)-8/5(火)

国立オリンピック記念青少年
総合センター(東京・代々木)

各種SNS

HP <https://zenmoken.com/%E7%AC%AC%E5%9B%9E%E3%BD%BDajemun/>

Instagram

<https://www.instagram.com/ajemun2025/>

Mail

ajemun2025@gmail.com

Youtube

<https://www.youtube.com/@ajemun2025>



QRコード

第9回AJEMUN 情報まとめサイト

<https://sites.google.com/view/ajemun-mogikokurenntaikai?usp=sharing>



目次

作者より	4
議題解説書について	5
議題解説書の位置付け	
議題解説書の扱い	
第1章 会議設定	6
第1節 議場設定	
第2節 成果文書について	
第2章 世界とAI技術の進展及びその影響	8
第1節 <u>AI技術の定義と歴史</u>	
第1項 AIとは	
第2項 AIの発展史	
第2節 <u>AIがもたらす社会的・経済的影響</u>	
第1項 生成AIがもたらすメリット	
第2項 生成AIがもたらすデメリット	
第3章 軍事用AIと国際安全保障	17
第1節 <u>軍事用AIの歴史</u>	
第2節 <u>軍事用AIの定義と現状</u>	
第1項 軍事用途におけるAIとは	
第2項 国際人道法との関係	
第3項 軍事用AIに関するリスク	
第3節 <u>軍事用AIを使った実例や軍事戦略の変化</u>	
第4章 国際的な規制議論と今後の方向性	27
第1節 <u>生成AIに対する各国の規制</u>	
第2節 <u>AIに関する国際的な議論の動向</u>	
第1項 これまでの国連の対応	
第2項 各国の法規制、ガイドライン等の整備動向	
第5章 論点解説	35
論点1 AIに対する規制	
論点2 AIの誤作動に対する責任の所在	
アウトオブオブアジェンダ解説	
参考	



作者より

この度は、AJEMUN第9回にご参加いただき、誠にありがとうございます。
皆さんと共に会議を創り上げられることを心から嬉しく思っています。

さて、今年度の議題は「AIと軍事」。この言葉を聞いて、皆さんはどのような印象を抱かれたでしょうか？

模擬国連の議題というと、時にスケールが大きく、何をどう考えれば良いのか見えづらいテーマも少なくありません。しかし今回の議題である「AI」は、もはや皆さんの日常と切り離せないほど身近な存在であり、世界でも急速に注目を集めているホットなテーマです。

例えば、勉強につまずいた時に質問に答えてくれるAI。予定管理や文章の添削、あるいは友人や家族にも話しづらい悩みをそっと聞いてくれるAI。また、皆さんの学校現場でも、AIを活用した教育支援の取り組みが進められているかもしれません。

そんな身近なAIですが、一方でそれが軍事の場で活用されるとしたら？

そのAIによって国際社会に甚大な被害がもたらされた時、私たちはどのような課題に直面するのでしょうか。そして、そのような事態を防ぐために、どんなルールや規制が必要とされるのでしょうか。

この問いに正面から向き合うのが、今回の議題です。

そして、これは少し意地悪な確認かもしれませんが、本会議に向けてのPPP作成や会議方針の決定などを、すべてAIに任せきりにしてしまっている大使はいませんか？

AIにヒントをもらうのはもちろん悪いことではありません。しかし、模擬国連における議論の場では、あくまで“大使の皆さん自身”の考えや判断が何よりも重要です。各国の大使として、国の立場や状況を理解し、自分の言葉で議論に向き合ってください。

私たちBG作成チームも、今回の執筆を通してAI技術の進化と多様な用途に触れ、その可能性の広さに何度も驚かされました。一方で、それが軍事と結びついたときに生じる倫理的な問題や国際的なルール作りの難しさにも直面し、深く考えさせられることもありました。

本会議を通して、皆さん一人ひとりが「AIとこれからどう関わっていくか」、そして「その技術が世界に及ぼす影響についてどのように考えるべきか」を、自国の立場を通して、見つめ直す機会になれば幸いです。

この議題には、正解があるわけではありません。しかしだからこそ、多様な意見や視点がぶつかり合う意味のある議論になると、私たちは信じています。

会議当日、皆さんがどのような視点を持って、どのような世界の未来を語るのか、その姿に出会えることを私たちは心から楽しみにしています。

(共立女子高等学校2年 北川眞子 鈴木花果 福岡加帆)



議題概説書について



議題解説書の位置付け

議題概説書は、今後皆さんが各国大使として準備を進めるにあたっての出発点として、会議の背景知識を共有するものである。本議題概説書を導入として、議題および担当する国への理解を深めていただきたい。リサーチを進める前に本議題概説書に目を通し、議題に関する歴史的背景や本議題の論点を整理していただければ幸いである。



議題概説書の扱い

1. 著作権は、全国教育模擬国連大会(AJEMUN)事務局に帰属する。
2. 本議題概説書を用いた学校間での練習会議を本大会終了まで禁ずる。
3. 大会終了後は、本議題解説書を以下の目的での使用を許可する。
 - 1) 学校内または学校間の練習会議での使用
 - 2) 校内での教育目的での使用

使用、引用する際は、出典を明記すること

出典の明記例:「AJEMUN 2025年度議題解説書 より」

その他の用途での利用を希望する場合は、事前にAJEMUN事務局へ確認を要する。

4. 改変・商用利用を禁ずる。



第1章 会議設定

本章では、会議の目的や設定、成果文書の扱いについて解説する。議題についての詳細なリサーチの出発点として、本議場は何を目的に開催されるのか、どのような役割を果たすべきなのか、大使の使命を明確にしていきたい。

第1節 議場設定

議場

今会議の議場は、第80回国際連合総会第1委員会（軍縮と国際安全保障）である。国際連合は1945年に設立された国際機関であり、現在193の加盟国で構成されている。世界のすべての国が集まり、共通の問題を議論することで、全人類に利益をもたらす共通の解決策を見つけることができる地球上唯一の場所である。

国連総会（UNGA）は、組織の主要な政策決定機関である。すべての加盟国で構成され、193ヶ国それぞれが平等な投票権を持っている。総会は毎年9月から12月までの定例会期で開催され、そのほかは必要に応じて開催される。

国連総会第1委員会（First Committee）は、国際連合総会の主要六委員会の一つであり、軍縮および国際安全保障を専門に扱う。Disarmament and International Security Committee（DISEC）とも呼ばれ、全193加盟国が参加し、毎年秋に開催される。委員会は核兵器、その他の大量破壊兵器、その他の軍縮措置および国際安全保障、軍縮機構の七つに焦点をあてる。

これらを通じて、国際的な安全保障体制の強化と恒久的平和の実現を目指す。

開催日時

本会議の開催日時は、2025年8月4日～5日と設定する。実際の第80回国際連合総会は2025年6月から2026年9月にかけて開催される予定だが、今会議では、第9回全国高校教育模擬国連大会（AJEMUN）当日の日程に開催される設定とする。よって会議では、AJEMUN前日までの情報を使用させていただいて構わない。

議論と論点

今会議の議題は「AIと軍事」である。

論点は、

- ①「AIの規制」
 - ②「AIの誤作動に対する責任帰属の明確化」の2つとする。
- ここでのAIとは、「生成AI」と「軍事用AI」のことを指す。
詳しい内容については次章以降を参照とする。

第2節 成果文書について

成果文書の扱い

第一委員会は、国連憲章第10条より、加盟国にとって法的拘束力を持たない「勧告（recommendation）」を作成する場である。

そのため、第一委員会で採択された文書は第一委員会の公式見解として国際社会に発信されるが、各国に強制力をもって実行を義務づけるものではない。

しかし、その「勧告」は、国際的な慣習法や規範の形成に影響を与えることもある。

したがって、本会議で採択される成果文書は、AIという強力な技術を適切に管理するための「国際的規範」としての役割を果たすことが期待される。AIは、その活用次第で人類に大きな利益も害ももたらし得るものであり、文書の意義を単なる声明にとどめることなく、今後の国際的な方向性を示す基盤とするべきである。そのことを踏まえ、各国大使の皆さんに有意義な会議をしていただきたい。

採択条件

国連憲章第18条第2項重要事項により、本会議に出席し投票する加盟国の3分の2以上の賛成によって行う。

上項目に記した通り、総会決議は原則として法的拘束力を持たないため、採択そのものが各国の行動を自動的に強制するわけではない。しかし、本会議の成果文書が、全世界のAI開発および事業に大きな影響を与えうることを留意していただき、決議案を具体的な行動が伴う、実効性のあるものでなければならない。

そのため、決議案は3分の2以上の賛成を目指すものではなく、「具体的行動を伴う実効性あるAIガバナンスの出発点」と位置づけるべきである。それらが各国大使の共通認識となるよう強く求める。



第2章 世界とAI技術の進展及びその影響



第1節 AI技術の定義と歴史



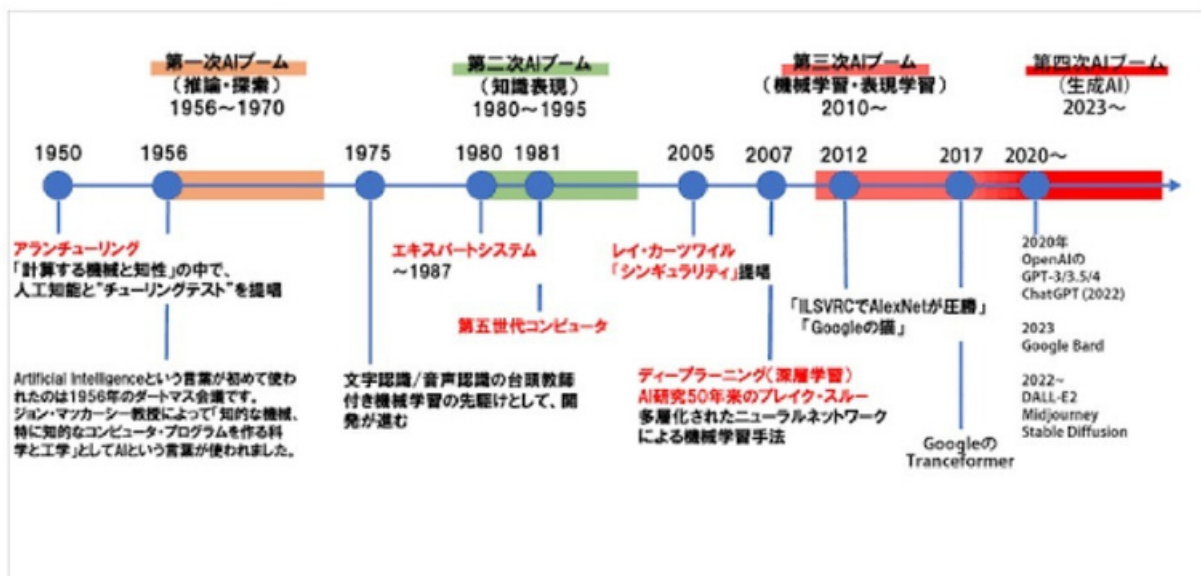
第1項 AIとは

AIとは、「Artificial Intelligence」の略称で、日本語では「人工知能」を意味する。AIは一般的に、人間の言葉の理解や認識、推論などの知的行動をコンピュータに行わせる技術を指す。人工知能の概念は1950年頃から存在していたが、「AI」という言葉は、1956年にアメリカのダートマス大学で行われた研究会において初めて公に用いられた。現在に存在する身近なAI搭載システムの例として、Siri、Alexaなどの音声アシスタント、ChatGPTなどのチャットbot、ルンバなどの掃除用ロボットなどがあるが、これらの多くは医療や企業などで活用されている。

生成AI (Generative AI) とは様々なコンテンツを新たに生み出すAIのことを指す。従来のAIが決められた行いを自動化するのに対し、生成AIはデータから学習した規則性や関係性を活用し、テキスト、画像、動画、音声など多岐にわたるコンテンツを新たに生成することができる。専門知識のない一般人でも比較的容易にコンテンツを作り出すことができるため、生成AIの更なる進化によって、生産性を向上させるほか、アートやエンタメなどの分野において表現の可能性を広げることができると言われている。



第2項 AIの発展史



↑ 資料 AI技術の年表

出典: <https://dxconso.com/about-ai/>

今までAIブームは4回(注参照)起こっており、そのブームごとにAIは発展していった。私たちが今いるこの時代は、第4次(注参照)AIブームの渦中である。

(注:現在のAIブームが「第3次」であるか「第4次」であるかについては、専門家の間でも見解が分かれており、明確に統一された定義は存在しない。これは、人工知能技術の発展過程における「ブーム」の区切り方や重視する技術的転換点に違いがあるためである。本議題解説書では、現在のAIブームを第4次として扱う。)

第一次AIブーム:探索と推論(1950～1970年代)

「AI」という言葉が初めて用いられた1950年代後半～1970年代に、欧米でAIの使用、開発が盛んになった。

「推論」や「探索」の技術を用いた、オセロやチェスなどのテーブルゲームをプレイできるAIが生まれる。

推論:AIが持っている知識やルールを使って、論理的に結論を導くプロセス。決まったルールに従って、確実に答えを導き出す技術。

探索:ルールだけでは答えが見つからないときに、たくさんの選択肢を試しながら最適な解決策を探すプロセス。迷いながら進むイメージで、答えを見つけ出す手法。

また、特定分野の「知識」を取り込んだコンピュータが、専門家のように推論を展開する「エキスパートシステム」が登場する。1966年には、今のチャットボットの元祖となる、初の対話型自然言語処理プログラム「ELIZA」が誕生した。

しかし、当時のAIは、現在のAIと比べて単純かつ小規模だった。ルールが明確に定められた単純な問題には対応できるものの、さまざまな要素や条件が複雑に絡み合う社会課題の解決には対応できず、第一次AIブームは終焉を迎えた。

第二次AIブーム:知識表現(1980～1990年代)

第二次AIブームは、コンピュータの活用が当たり前になっていく1980～1990年代に起こった。第一次AIブームで登場したエキスパートシステムがビジネス(株価予測)や、医療(病理診断)など、幅広い分野で実用化されていった。

しかし、AIの判断の精度を高めるために必要な「知識(=コンピュータが理解できる形で記述された膨大な量のデータ)」を、人間が用意しなければならなかった。また、「暗黙知」と呼ばれる言語化しにくい知識のデータ化が困難だったこと、例外的なルールや処理に対応しきれないことなどから、第二次AIブームも次第に勢いを失っていった。

第三次AIブーム:機械学習(2000年代～)

2000年代になるとコンピュータの性能が飛躍的に向上し、現在まで続く第三次AIブームが起こった。大きな飛躍の契機となったのが、コンピュータが大量のデータからルールやパターンを発見し、自動でものごとを学習する「機械学習」の登場だ。その後、2006年にはこの技術を発展させ、さらに複雑な判断が可能となる「深層学習（ディープラーニング）」が実用化され、歴史的なブレークスルーとなった。第2次AIブームの課題となっていた、データを用意する手間が劇的に軽減されたことで、さまざまな分野に導入され実用化が進んでいる。

第四次AIブーム：生成AI

時代の最先端として注目を集めているのが生成AIである。入力したテキストの指示に基づいて画像データを出力するAI（Midjourneyなど）や、人間と話すような感覚で文章生成や情報収集を行ってくれるAI（ChatGPTなど）などが大きな話題となっている。倫理的・法的なリスクが指摘されている面はあるものの、今後、さらに多くの場面で活用されることが予測されている。



第2節 生成AIがもたらす社会的・経済的影響



第1項 生成AIがもたらすメリット

教育

AIは教育分野において、個別最適化学習から教材作成・授業支援、評価の自動化、遠隔学習の推進まで多岐にわたるメリットをもたらしている。

①個人に合わせた教育

AIによって、すべての人々が一人ひとりに合った教育を受けられる、個別最適化学習ができるようになった。個人の理解度や進捗に合わせて教材やペースを自動調整することをチューターなどがいなくても可能になる。

②自動採点 & リアルタイムフィードバック

学習履歴や行動データを解析し、早期に「つまずきパターン」を検知することができる。介入すべきポイントを教師に通知して、手遅れを防ぐことも可能である。

③学べる場が少ない社会人にも

社会人向けリスニングや語学学習にも、24時間365日、AIが個人の成長を促進する。

また、教師にとってのメリットとして、業務の効率化がある。ルーティン業務（出席管理・成績集計・教材作成の補助）をAIに任せることで、教師は本質的な指導や対話に集中することが出来る。ゲーミフィケーションやチャットボットによる対話型学習で、生徒の興味を引き出し、勉強のモチベーションを高めるなどの活用の仕方考えられる。

医療

AIは医療分野において、診断精度の向上から業務効率化、遠隔医療、創薬、個別化治療まで多岐にわたるメリットをもたらしている。

①診断精度の向上&予測分析による早期介入

生成AIは大量の医療データを迅速に解析し、従来の方法では見落としがちな微細なパターンを検出できる。たとえば、AIによる画像解析は、X線やCT、MRIなどでの異常検知を支援し、早期発見率を高めることで誤診リスクを低減することができる。また、AIによる予測分析は、患者の診療履歴や検査データから疾患発症リスクを予測し、未然に対策を打つことを可能にする。入院ピークやベッド需要を予測してリソース配分を最適化したり、再入院リスクの高い患者を特定して退院後のフォローアップを強化したりと、医療の質と効率を同時に向上できる。これにより、重症化や緊急搬送の回避にもつながる。

②創薬プロセスの加速

AIは、分子構造の予測や候補化合物のスクリーニングを高速化し、従来数年を要した創薬フェーズを大幅に短縮させることができる。機械学習モデルは作る薬の有効性や安全性を予測し、不適切な創薬過程を早期に除外することで、医薬品の研究開発コストを抑制する。また、AIは新たな標的分子の発見や最適化にも貢献し、より効果的で副作用の少ない治療薬の開発を後押しする。

③遠隔医療とアクセス拡大

AI搭載の遠隔診療プラットフォームは、都市部・過疎地を問わず専門医の診断を届ける仕組みを提供することができる。画像診断支援やバーチャルアバターによるトリアージ、チャットボットによる初期間診で、患者は自宅から高品質な医療サービスを受けることができる。これにより、患者の通院負担や待ち時間を削減し、医療格差の是正にもつながる。

以上のように、医療体制が遅れている国家にとっては、AIは医療体制を整えてくれるものとなる。

労働

労働分野においても、AIは作業の自動化による業務効率の向上から、人手不足の解消、労働環境の改善など、幅広い利点をもたらしている。

①効率化

AIは、定型的な業務や繰り返しの作業を自動化し、従業員はより戦略的・創造的なタスクに集中することができる。AIアシスタントを導入したチームは、平均で15%の生産性向上を達成するなどのデータもある。このような事例は、とくに経験の浅いメンバーとの比較で顕著に現れる。AIは、従業員の知識ギャップを埋める“仮想チームメイト”として機能することが可能と言うことだ。この点は、労働環境の技術格差が大きい地域や、雇用体制が整っていない地域、労働環境が整っていない地域に対して有効だろう。

②コスト削減

AIは、生産性向上とエラー削減により、組織全体の運営コストを抑制させることができる。また、製品・サービスの効率的な生産で利益率が向上する。これらは、高齢化や労働力不足に直面する国、財政制約に厳しい新興国に対し有効だろう。

③安全性の強化

センサーやカメラと連携し、リアルタイムで危険な兆候を検出することができる。また、事故や労働災害を未然に防止する仕組みを提供する。このようなメリットは、テロなど職場や施設の安全面の不安が多い地域で有効だろう。

④ワークライフバランスの向上

AIの使用と業務の自動化により、従業員の89%が仕事の満足度が上がり、91%がワークライフバランスの改善を実感できる。柔軟な働き方を支援し、従業員エンゲージメントを高めることができる。これらは、長時間労働が横行している国家や人材流出に悩む国々、イノベーション競争が激しい先進国に対して有効だろう。

民間において

生成AIをメンタルケアマネージャーとして活用してる人々が多い。AIを思いのほけ口や相談相手とすることで、日々の生活を楽しめることができる人々が増えるだろう。これらは、自殺率の高い国家や幸福度が低いとされる国において有効だろう。



第2項 生成AIがもたらすデメリット

雇用喪失・職業の消失

AIによる業務自動化により、単純作業や定型業務がAIに置き換えられ、多くの職種で雇用が減少するリスクがある。これらは、失業率が高い国に対して大きく見られるデメリットだろう。

アルゴリズムバイアス・差別

アルゴリズムバイアスとは、AIが特定の属性（人種、性別、年齢、宗教など）を持つグループに対して、不公平な結果を生み出す現象のことである。例えば、採用選考や融資判断などで特定の性別・人種に対して不利な判断を下すことなどがある。これらは、データ偏在が顕著な途上国、AIインフラの整備が不十分な国、差別が生じやすい多民族、格差社会の国家に対して大きくみられるデメリットだろう。

プライバシーの侵害・ディープフェイク・偽情報拡散

AIの利用によって、個人のプライバシーは多方面から脅かされる。駅や街頭に設置されたAIカメラによる顔認識で、個人の移動や感情（表情から推定される心理状態）も常時追跡されるようになり、監視社会化が進行している。さらに、AIによる個人の行動履歴や属性から、自動的に評価や判断を下す仕組みである自動プロファイリングが融資や採用判断に使われ、不透明な基準で特定の属性を不利に扱う差別も生まれている。また、本人そっくりの偽映像や音声が生動生成される「ディープフェイク」が近年、社会的関心を集めており名誉毀損や詐欺に悪用される事例も増加している。

これらの問題は、特にプライバシー保護に関する法制が整っていない新興国において深刻である。民主的な監視機構が不十分な途上国では、市民が知らないうちに国家によってデータを奪われ、言論の自由が萎縮させられる可能性が高くなる。そのような国家が現れた際の国際的な制裁も必要である。

ブラックボックス問題

AIの判断プロセスが複雑であるために、人間にはその根拠や意図が見えづらく、「なぜその結論に至ったか」が説明できないことがある。そのため生成AIにおいて構造が複雑で判断根拠が見えないことを「ブラックボックス問題」と言う。AIの出力に誤りや不具合が生じた場合、つまり“誤作動”を起こした際の原因解明が困難で、生成AIそのものの信頼性と安全性が損なわれる。例えば、医療や金融のような説明責任が求められる分野で、AIの判断理由が不明なことは致命的である。設計や学習データが非公開である場合、AIの判断に差別を含まないかを調査するバイアス監査を行うことができず、差別された人にとっての不公平な結果が見逃されやすい。結果として、AIが誤作動をした責任が曖昧になり、誤作動によって被害を受けた人が救済されにくい構造が生まれている。また、この説明責任の欠如が、不正やバイアスの温床となる可能性がある。

社会的不平等の拡大

高価なAI技術を保有する大企業や先進国が、研究開発資金やデータアクセスを独占することで、最新のAIを活用できる組織と、導入コストの捻出が困難な中小企業・発展途上国との間に大きな技術格差が生じる。

また、優秀なエンジニアや研究者は高待遇を提示する大手企業へ流出し、人材育成の余力を持たない地域では高度なスキルが定着しにくい。その結果、生産性や収益は極少数に集中し、AIの影響力を背景とした政策決定や社会制度の形成においても意見の偏在が発生する。こうした状況が続くと、教育・医療・交通などの社会インフラへのアクセスにおいて、「AI技術を有する者」と「有しない者」の間で格差が拡大し、デジタルデバインドが一層深刻化する恐れがある。

認知機能の低下・過度依存

AIに過度に依存すると、人間が自分の頭で考えたり覚えたりする場面が減るため、思考力や記憶力が衰える可能性があると言われている。たとえば、検索エンジンや生成AIにすぐに答えを求める習慣が身につくと、自分で情報を整理したり、問題を深く掘り下げて考えたりする機会が少なくなる。また、必要な情報をすぐにAIが提示してくれるため、「覚える必要がない」と脳が判断し、記憶力の低下にもつながると言われている。こうした傾向は、特に学習や仕事の場面で顕著であり、AIが便利である一方で、人間の認知能力の一部を徐々に退化させるリスクがある。このように、AIの活用が進む中で、「考える力」を維持・育成する意識がこれまで以上に重要になってくる。

セキュリティリスク・軍事的利用

AIを使ったサイバー攻撃（フィッシングメール自動生成やマルウェアの高度化）や、テロリストがAIを利用して、無人ドローンの自動操縦や顔認識による標的攻撃など、より正確で効率的な兵器、攻撃方法を行う危険性がある。また、AIによって爆弾を爆発するタイミングや場所を最適化することも可能となり、被害が拡大するおそれがある。さらに、AIを使ったサイバー攻撃でインフラや通信網が狙われるリスクも高まっている。こうした技術の悪用は、従来のテロよりも脅威を増大させる可能性がある。また、AI搭載ドローンやロボット兵器の自律的な判断が誤作動を起こした場合、責任所在や倫理的問題が深刻化する。この問題は論点2で詳しく扱っていく。

ここではメリットとデメリットとしてそれぞれの要素を挙げたが、これらは国家によって真逆の意味合いを持つ場合もある。例えば、AIによる業務の効率化は、ある国では労働負担の軽減として歓迎されるが、別の国では失業の深刻化につながる可能性がある。つまり、メリットかデメリットかという評価自体が国の経済状況や政策、社会構造によって大きく変わり得る。

また、AIの判断過程が不透明な「ブラックボックス問題」なども、一部の権威主義国家において、監視強化や世論操作などの手段として悪用され、人権侵害の実態を覆い隠す温床となる恐れがある。

したがって、自国の実情を踏まえ、どのような視点でこの問題を評価し、どのような原則を支持するかは、各国大使のリサーチと判断にかかっている。慎重かつ主体的な見極めが求められる。



第3章 軍事用AIと国際安全保障



第1節 軍事用AIの歴史

軍事革命

人類の歴史において、戦争の様相を根底から変えた軍事的革命を「軍事革命」という。ここでは火薬、核、AIという3つの技術革新に基づく軍事革命を順に概観する。

第一次軍事革命

16世紀から17世紀にかけて、火薬兵器（マスケット銃・大砲）の普及が戦争の様相を劇的に変化させた。それまで主流だった騎士や弓兵は姿を消し、訓練された歩兵による火力重視の戦術（線形戦術）が登場する。

大砲の威力に対応して「星形要塞」のような新たな防衛構造も開発された。この技術革新は、単なる戦術変更にとどまらず、戦争を長期的・制度的に遂行する常備軍と官僚制度の発展を促し、結果として中央集権国家の形成にもつながりました。第一次軍事革命は「国家による戦争」の始まりを告げた転換点になった。



星型の要塞都市

引用:算数・数学な日々 - 大日本図書

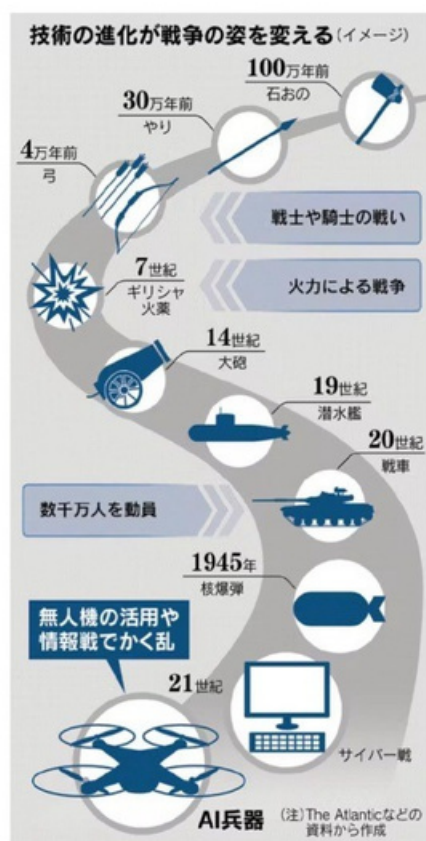
https://www.dainippon-tosho.co.jp/mathful_diary/2023/0707.html

第二次軍事革命

20世紀に入ると、産業革命による大量生産の進展とともに、第一次・第二次世界大戦で機関銃、戦車、航空機などの新兵器が登場し、総力戦が可能になった。しかし、軍事の構造そのものを根底から変えたのは、1945年に登場した核兵器だった。核兵器は、一発で都市を壊滅させる圧倒的破壊力があつた。もし核兵器を使って戦争になれば、お互いに相手国を完全に破壊してしまうことが確実だという相互確証破壊（MAD）に支配されるようになった。そのことから「戦争を始めるかどうか」という決断が重大な問題となり、戦争の目的・抑止・外交までもが再構築された。これらのことから第二次軍事革命は核革命とも呼ばれる。

第三次軍事革命

現在進行中の第三次軍事革命は、AI(人工知能)と自律兵器の登場によって進んでいる。これまで兵器の使用には必ず人間の判断が介在していたが、AI兵器は敵の識別・判断・攻撃をすべて自律的に行う能力を持ち始めている。さらに、ドローンやサイバー兵器、精密誘導ミサイルとの統合により、戦争は「物理的な戦場」だけでなく、情報・心理・宇宙・サイバー空間へと拡張している。AIの導入により、戦争のスピードは人間の反応速度を超え、人間が指揮権を失う可能性も懸念されている。このようにAI革命は、単なる技術革新ではなく、「誰が敵か、いつ攻撃するか」という戦争の根本ルールを書き換える革命として、核兵器以来の決定的変化とみなされている。



↑ 資料 兵器の年表
出典: 米軍がAIに負ける日

<https://www.nikkei.com/article/DGXMZO43394860V00C19A4TCR000/>



第2節 軍事用AIの定義と現状



第1項 軍事用途におけるAIとは

AIの軍事への応用分野は、情報収集・分析、意思決定支援、自律型システム、サイバー攻撃、ロジスティクスなど多岐にわたる。近年、特に議論が活発に進められているのが、自律性を持つ軍事用AIシステムである。

無人兵器の分類

無人兵器は、「遠隔操作型」、「半自律型」、「完全自律型」の3つに分類することができる。

「遠隔操作型」は、目標の選択、爆弾の投下、突入などすべての操作を人間が行う。たとえば、システムによって目標の捕捉(目標を感知、認識すること)や追従(目標の進路に合わせて追尾、追跡すること)を自動化しても、弾薬の投射や突入(自爆)は人間の直接的な操作で行われる。

「半自律型」は、人間が攻撃目標を設定し、ロボットは設定された特定の目標を探索・追尾し、攻撃できる時期を見計らって爆薬を投射もしくは突入して自爆する。爆発の引き金は、

①機械が人間の最終判断を受けて作動する場合

②機械が独自に判断して作動する場合

後者の場合は「完全自律型」に近い性質を持つとされる。

「完全自律型」は、目標の決定から攻撃までをシステムが行う。

人間が設定した「敵兵士」や「戦車」、「弾薬庫」などの情報をもとに、戦場で目標を探して攻撃する。「半自律型」の②と「完全自律型」の違いとして、②は人間が最初にターゲットや目標エリアを選定し、機械がその枠内で独自判断することが挙げられる。対して、「完全自律型」は先述の通り、目標設定から攻撃まですべてAIが行う。



↑ 資料 LAWSの分類: 遠隔操作型、半自律型、自律側の違い

出典: 恐怖のAI兵器「LAWS(自律型致死兵器システム)」とは? 使用規制が進まない複雑な事情

<https://www.sbbt.jp/article/cont1/152252>

無人兵器の利点

人間の戦闘員の代わりに、Dull(退屈)、Dirty(汚い)、Dangerous(危険)、Deep(深い)の4つの「D任務」を遂行することができる。またがない分、人間より正確に人命に影響を与えるような殺傷力の高い兵器致死性武力の行使が可能であり、人間特有の恐怖、疲労、復讐心といった人間の感情や体調から生じる判断のミス回避できる。一部の戦争犯罪は仲間の死に対する復讐を求めることで起こる。そこで、感情がなかったり、仲間に対して愛着を覚えなかったりする無人兵器を使用することは、民間人の危険を軽減できると考えられている。

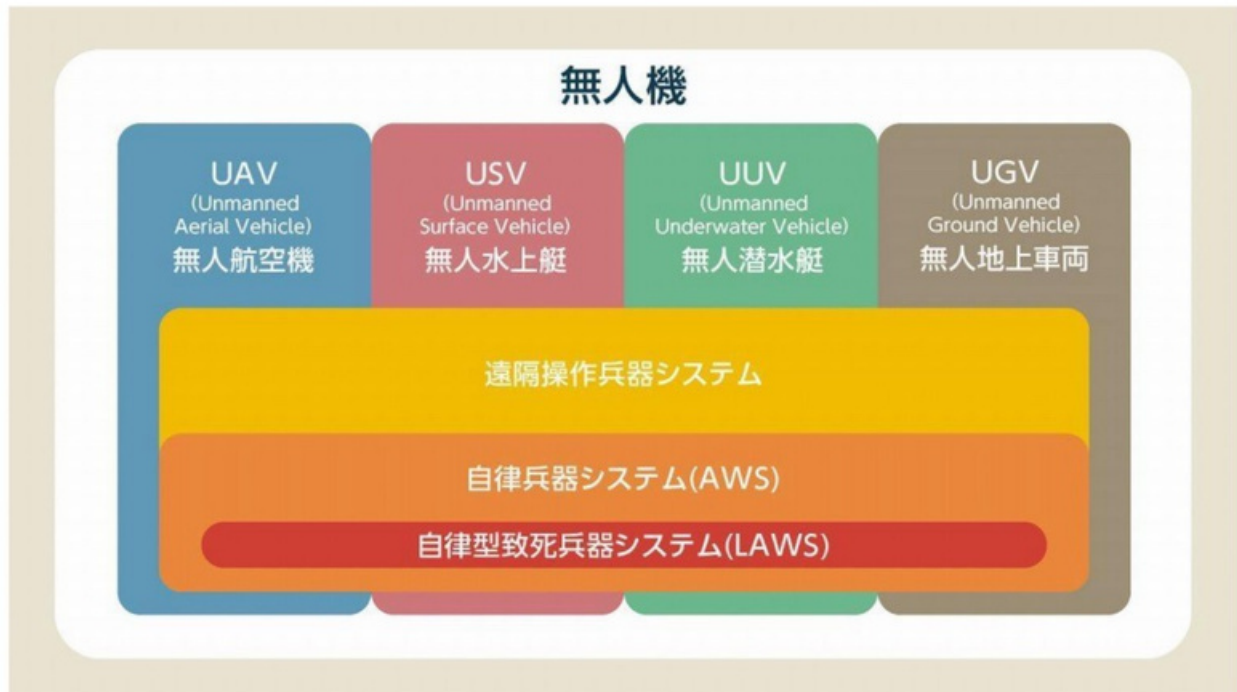
No	区 分	内 容
①	Dull (退屈)	長時間監視任務
②	Dirty (汚い)	化学・生物・放射性物質・核(CBRN)環境下での任務
③	Dangerous (危険)	対空脅威の高い領域での任務
④	Deep (深い)	敵領域内での長距離移動を要する任務

また、自律性を有する無人兵器の利点は、システムの運用に必要な人的作業負担を削減することは元より、最も重要な利点はスピードである。

自律性を持つことで、観察(Observe: 目標搜索)、情勢への適応(Orient: 目標探知)、意思決定(Decide: 探知目標への対応決定)、行動(Act: 探知目標への攻撃)の、いわゆるOODAループを

早く実行できる。また、人間より複雑さや量、スピードにおいて、高度にデータを収集、照合、分析できるようになる。これにより、意思決定サイクルの質とスピードが大幅に向上することが指摘されている。

LAWSとは



↑資料 LAWSの位置付け

出典: 恐怖のAI兵器「LAWS(自律型致死兵器システム)」とは? 使用規制が進まない複雑な事情

<https://www.sbb.it.jp/article/cont1/152252>

人間による遠隔操作を必要としない兵器のことを、「自律型の兵器(AWS: Autonomous Weapon Systems)」という。そして、AWSの中でも殺傷能力を持つ兵器を「自律型致死兵器システム(LAWS: Lethal Autonomous Weapons Systems)」と呼ぶ。

国際法における定義については議論が進められている段階だが、基本的には「攻撃目標の選定」から「攻撃の実行」までのプロセスに、人間がまったく介入せずに動作することができる兵器についてはLAWSであると考えられている。また、敵の認知や攻撃の判断など膨大な情報処理をAIが主体的に行うことで、兵器が「自律」的に相手を「致死」させるかどうかを判断するため、LAWSは兵器単体ではなく、人間の判断を介在させない攻撃を可能にするシステム全体を指すとされることもある。

先述の通り、LAWSの定義については、国際社会で議論が行われており、未だ定まっていない。国によって細かな定義が異なるため各自リサーチを行なってほしい。



第2項 国際人道法との関係

「国際人道法」とは

「国際人権法」と同様、「国際人道法」という名の条約や法律があるわけではなく、戦闘の方法や手段を具体的に制限・禁止し、戦時下で保護する対象を明文化した条約や慣習法などさまざまな国際的なルールの総称である。

国際人道法の主な原則は、「区別」、「均衡性」、「予防」の3つとされている。また、LAWSを含むすべての兵器はこれらの原則に従う必要がある。

①「区別の原則」

攻撃側は常に戦闘員と民間人、軍事目標と民用物（住宅や民間インフラ、環境）を区別しなければならないことを意味する。特定の軍事目標に向けられない攻撃、いわゆる無差別攻撃も禁止している。

②「均衡性の原則」

攻撃により民間人が巻き込まれて死傷したり、民用物に損害を与えたりすることが予想され、そうした犠牲が具体的かつ直接的な軍事的利益と照らし合わせて過大となる場合に、攻撃自体が禁止されるというもの。軍事作戦から予測される被害とは、直接的な影響だけでなく、インフラや環境への第二次、第三次的な影響も含まれる。

③「予防の原則」

民間人や民用物への危害を回避するか、少なくとも最小限に抑えることが可能な場合、紛争当事者は予防措置を講じなければならないというもの。上述の区別と均衡性の義務を尊重するための原則である。



第3項 軍事用AIに関するリスク

天候トラブルなどによる制御不能

遠隔操作されているドローンや無人兵器システムは、急激な気象変化（突風、豪雨、砂嵐など）に弱く、センサーの誤作動や通信の途絶が起こる可能性がある。こうした状況では、AIが誤ってターゲットを識別したり、操縦不能に陥ったりすることが考えられ、制御の喪失は重大な事故や誤爆につながる可能性がある。

過去に、以下のような実例があった。

Heron偵察ドローンが突風による進入失敗と地上ヘリ衝突（2025年3月17日、韓国）

突発的なアップドラフト（上昇気流）により機体が跳ね上がり、さらに横風と重なって滑走路を外れて駐機中のヘリに激突。軽火災が発生し両機とも損失 京畿道楊州市にある陸軍所属の航空大隊で、着陸を試みた無人偵察機1機が、突発的な上昇気流により機体が跳ね上がった。さらに横風と重なって滑走路を外れてしまい、飛行場に駐機していた多用途機動ヘリコプターと衝突した。軽火災が発生し両機とも損失。被害額は200億ウォン（日本円では約20億円）を超えた。



実際の様子 写真:NEWSISより

誤認識

AIは画像認識やパターン分析に長けている一方で、人間のような分別や常識は持ち合わせていない。

例えば、銃を持っている人物がいるとして、彼らが「狩りをしようとしているのか」、はたまた「戦闘行為をしようとしているのか」など、

具体的な状況を識別できない。

そして、これらの誤認識によって、戦争とは無関係な民間人が多くの被害を被る可能性がある。従来の兵器は人間が操作するため、誤った判断や意図しない行為が行われた場合、その責任は操作者に帰属した。しかし、軍事で自律型兵器を利用する場合、最終的な判断はAIが行うため、責任の所在が曖昧になる。そのため、AIに責任を問うことはできず、開発者や製造者、運用者など、様々な主体が責任の一端を負う可能性が考えられるだろう。この責任の所在の不明確さは、国際法の枠組みにおいても大きな問題となる。既存の国際法は、人間による行為を前提としており、自律型兵器のような非人間的な行為に対する規定が十分ではないのだ。

実際の例として以下のような事件があった。

①カブール車両誤爆（2021年、アフガニスタン）

アフガニスタン首都カブールで、米軍ドローン(遠隔操作型)がISIS-K候補と判定した車両を誤認し射撃。結果、民間人10名(うち子ども7名)が死亡。

②STM Kargu-2(2020年、リビア)

リビア第2内戦末期にて、ドローン自身が退却中の敵部隊を目標を設定し、自律的に攻撃した。UNエキスパート報告では、複数の戦闘員が戦死した可能性が示唆されているが、正確な死者数は不明。

特徴として、半自律型の②に分類される。すなわちオペレーターが派遣や導入を指示し、攻撃判断はドローン側が独自判断する構成である。

影響:2021年の国連報告書では、自律判断による初の実戦使用例として記録されている



↑ 世界初の殺人ロボットと報道されたSTM Kargu-2
写真:東京新聞より



第3節 軍事用AIを使った実例や軍事戦略の変化

軍事用AI導入による戦争被害

2022年2月24日、ロシアがウクライナへの侵攻を開始した。ウクライナは、ロシアとの軍事力や資源、人員の差を埋めるために、AI兵器の開発に積極的に取り組んでいる。遠隔操作型の無人機で偵察や監視、攻撃を行い、現場との連携体制を築いてきた。しかし、ロシアが対抗策として電波妨害を行ったことで、無人機の通信が途絶えて制御不能になるケースが増えた。そのため、ウクライナはAIを搭載した無人機を戦地に投入し始めている。これらは遠隔操作を必要としないため、電波妨害の影響を受けずに自律飛行ができる。ウクライナによる2023年10月の発表によると、このようなAIを搭載し、自動で目標の検知と追跡を行う無人機をすでに2000機配備している。

中東では、2023年10月7日にイスラエルとハマスらがガザ地区に拠点を置く武装勢力とのあいだで衝突が起こった。イスラエル軍は、AIを使った兵器やシステムを導入していると発表をしている。目標の捕捉(目標を探知、認識すること)における誤認率が10%もある状況下で、AIシステムを操縦するオペレーターたちは目標の確実性を検証することも無く、AIに攻撃させる判断をしている。本来目標に定められていないはずの民間人の巻き添えも報じられている。

これらの事例からも、戦場に投下される兵士の数を減らすことは出来ても、民間人の犠牲は減ることがないように思える。

軍事戦略の変化

AIの軍事への応用分野は、情報収集・分析、意思決定支援、自律型システム、サイバー攻撃、ロジスティクス(作戦に必要な物資の補給や整備、連絡などにあたる行動)など多岐にわたる。これまで、人間の意思決定や操作に頼っていた戦闘行動が、AIによって部分的または完全に自動化されることで、迅速かつ正確な判断が可能となる。例えば、自律型ドローンや自動化された監視システムがAIを搭載し、リアルタイムで目標を特定し、攻撃を実行に移す。これは、戦場における反応速度を飛躍的に向上させ、犠牲者の減少につながる可能性がある。一方で、AIが人間の本来意図した目的、指示、規範から外れる「誤作動」や人間の指示に反する判断を下す「誤判断」が深刻な被害を引き起こすリスクも無視できない。無人兵器が誤って市民を攻撃するような事態が発生すれば、国際法違反や人権侵害につながる可能性がある。また、AIを利用した戦争が拡大することで、国家間の緊張が高まり、AI技術が軍事競争を激化させるリスクも存在する。また、軍事用AIの導入により、兵士を戦場に送らず人的被害を大幅に抑制できる。だが、AIが誤って一般市民を敵と認識すると自動追跡が行われ、巻き添え被害がおきるリスクがある

そして、攻撃の判断を下すオペレーターは、ボタン1つで画面越しに簡単に人を殺すことができてしまい、基本的人権が危険に晒されてしまう可能性がある。

一方で、軍事用AIは今世紀の重大な抑止力にもなると言える。現在の兵力差を技術開発により縮めることもできる。AIの迅速で精密な判断は迅速な対応につながり、

無人兵器や自立型システムが配備されることで敵は安易に攻撃を仕掛けにくくなることやAIの反応を敵が完全に予測できないことで攻撃を思いとどまらせる可能性がある。そして、AIを人的戦力に代替したり、高精度な情報処理と戦術支援をおこなえたり、高価な戦闘機や空母がなくても、AIドローン群やサイバーAIによって先進国とある程度渡り合えるため現在の人的兵力差を縮めることができる。



第4章 国際的な規制議論と今後の方向性



第1節 生成AIに対する各国の規制

身近な生成AIの代表例である、ChatGPTに対する各国の規制をいくつか例示したいと思う。

★ Chat GPT 規制国の代表例

・中国

中国政府は本土内でのChatGPTへのアクセスをブロックしている。しかし、中国独自でChatGPTの代わりになるアプリ(バイドゥ社の文心一言など)が作成されており、それらは使用可能である。

・北朝鮮

北朝鮮では、政府がインターネットアクセスを厳しく制限しており、一般市民がChatGPTを含む外部のオンラインサービスを利用することはほぼ不可能である。

・イラン

イラン政府は、インターネット上の情報を厳しく監視・制限しており、ChatGPTへのアクセスもブロックしている。

・ロシア

ロシアでも、政府が特定のオンラインサービスやウェブサイトへのアクセスを制限しており、ChatGPTもその対象となっている。

★ Chat GPT の規制に対して肯定的な国の代表例

・イタリア

2023年3月、イタリアのデータ保護局は、プライバシー上の懸念からChatGPTへのアクセスを一時的にブロックした。その処置を受けopenAI社はchatGPTの安全対策の仕組みを公表した。そこでイタリアは、openAIが個人データの取り扱いを改善したと判断し、2023年4月28日に使用禁止を解除した。

・EU諸国

イタリアの措置を受け、EU諸国でもChatGPTのブロックを検討している。



第2節 AIに関する国際的な議論の動向



第1項 これまでの国連の対応

特定通常兵器兵器使用禁止条約(CCW条約)

1980年10月10日に採択、1983年12月2日発効、ジュネーブで署名された。これは、通常兵器(核兵器以外の兵器)の中でも、とくに非人道的被害をもたらすものを対象に、その使用禁止や制限を国際的に議論した。禁止・制限された主な兵器には、敵と民間人を区別せずに攻撃してしまう非識別兵器、焼夷兵器、地雷、即席爆発装置、レーザー兵器、戦争残存物がある。ただ、この条約は、核兵器や通常の銃器など多くの兵器が対象外であり、各議定書への参加は選択的批准であるほか、条約の執行機関がないため、強制力が弱いものになっている。

2013年に国連の「特定通常兵器使用禁止制限条約(CCW)」の枠組みでLAWSIに関する議論が公式に開始された。CCWは、非人道的な通常兵器(例:地雷、焼夷弾、レーザー兵器など)の使用を規制・禁止するための条約で、1980年に採択されている。この年から、各国政府や専門家、市民団体、技術者たちがAI兵器のリスクや規制の可能性について話し合う場として、CCWの中の会合が開かれるようになる。

2014年には政府間専門家会合(GGE: Group of Governmental Experts)が発足し、LAWSIに関する定義の検討、技術的な現状、国際人道法との関係性、倫理的・法的懸念といったテーマについて本格的な討議が始まった。その後の会合では、「人間の関与(meaningful human control)」という考え方が重要な論点として定着していった。これは、AIが殺傷判断を行う場合でも、一定の段階で人間が責任ある意思決定に関与する必要があるという原則である。

2016年～20年この時期には、「意味ある人間の関与(MHC)」という概念が重要なキーワードとして定着する。これは、LAWSの運用において人間が最終的な判断を下す必要があるとする立場で、多くの国がこれを支持した。

2021年、GGEは各国が共有する懸念や原則についての共通理解を文章化し、LAWSIに対する国際的な規範の方向性を示した包括的な報告書を提出した。詳しくは下記に記す。

2022～23年頃から、オーストリア、ニュージーランド、メキシコ、ブラジルなどの中小国グループが、LAWSIに対する新たな国際条約の交渉開始を求める動きを強めた。また、国際NGOや市民団体(特に「Stop Killer Robots」キャンペーン)も、国連でのプレッシャーを強め、「人間が介在しない殺人兵器は人道に反する」として全面禁止を求めている。

2024～25年アメリカ、中国、ロシアなどの大国は、「AI技術はまだ成熟しておらず、フル自律兵器の実用化は現実的でない」「研究の自由を制限すべきではない」などとして、全面的な禁止や新条約の策定には否定的・慎重である。そのため、国連CCWの場では、一部の国が「新たな法的拘束力ある規制」を求め、他の国は「既存の国際人道法で十分対応可能」という立場をとっており、各国の立場の隔たりが大きく、法的ルール作りには至っていない。

GGE報告書の概要と意義

2021年、GGEは各国が共有する懸念や原則についての共通理解を文章化し、LAWSに対する国際的な規範の方向性を示した包括的な報告書を提出した。以下に主な内容を詳述する。

1「意味ある人間の関与(Meaningful Human Control)」の確認

報告書では、兵器の運用に際して「意味ある人間の関与」が不可欠であるとされ、兵器が完全自律的に殺傷行為を行うことに対する懸念が共有された。目標の識別・攻撃判断・実行の各段階において、人間が関与することが国際人道法(IHL)を守る上で必要であると強調された。

2 国際人道法との整合性

GGEは、LAWSの設計・開発・使用が国際人道法に従う必要があることを明確にした。特に、区別(combatantと非combatantの識別)、比例(攻撃の被害と軍事的利益の均衡)、必要性(軍事目標への限定的使用)といったIHLの基本原則の順守が求められる。

3 責任の所在の問題

自律兵器の運用結果について、開発者・操作者・指揮官・国家のいずれが責任を負うのかという課題についても報告書は言及している。誤作動や予測不能な挙動による被害が発生した場合、明確な責任の所在がないことは、国際法の観点から重大な問題であると指摘された。

4 技術的發展と民生用途との重なり

報告書では、AIやセンサー技術などの多くが民生用途と軍事用途の両面を持つことから、研究・開発自体を規制することには慎重な姿勢が示された。そのため、使用の段階における明確な規範の確立が現実的な対応策とされた。

5 法的拘束力の欠如

2021年報告書に含まれる合意内容は、政治的・道義的な意味はあるものの、国際条約としての法的拘束力は持たない。各国の自発的な取り組みに依存しているため、強制力のある国際的な規制にはつながっていない。

国連教育科学文化機関AI倫理勧告

2011年11月24日に採択された。世界で初めて採択されたAI論理に関する包括的な国際基準。人間の尊厳、基本的人権、社会的公正、包括性、持続可能性、ジェンダー平等などの倫理的価値観を中核に捉えている。とくに発展途上国、マイノリティ、女性や子供といった社会的に弱い立場にある人々への配慮を重視している。

OECD AI原則

2019年5月22日にOECD(経済協力開発機構)が採択した原則。これは、AIの開発と利用において、国際的に共有できる価値と行動方針を示した世界初の政府間合意である。法的規制ではな

く柔軟なガイドラインとして機能し、各国が自国の状況に応じたAI戦略を策定・調節できるよう促す。

AIの発展による経済成長を推進しつつ、AI使用によって犯される危険性がある倫理的な観点にも着目した原則である。

GP AI (AIに関するグローバルパートナーシップ)

2020年、人間中心の考え方に立ち、「責任あるAI」の開発・利用を実現するために、OECDとG7の共同声明により創設された組織。

この組織は、OECDが事務局を務め、価値観を共有する政府、国際機関、産業界、有識者等からなる官民国際連携組織で、現在29か国が参加している。

GP AIには、「責任あるAI」「データ・ガバナンス」「仕事の未来」「イノベーションと商業化」という4つの研究部会が設置されており、専門家による議論と実践的な調査が実施されている。

GP AIの年次サミットである「GP AIサミット 2023」においては、新たなGP AI専門家支援センターである、GP AI東京専門家支援センターの立ち上げが承認された。

同センターでは、生成AIに関する調査・分析等のプロジェクトを先行的に実施する予定となっている。

ベレン宣言 (Belén Communiqué)

2023年2月24日、コスタリカのヘレディアにおいて、ラテンアメリカおよびカリブ海諸国 (CELAC) 33か国は、致死性自律兵器システム (LAWS) に関する重要な共同宣言、「ベレン宣言 (Belén Communiqué)」を採択した。この宣言は、同地域の国々が一致して、国際的な法的枠組みを通じて自律兵器を規制・禁止すべきだと訴えた。また、兵器システムの運用において、生命や死を左右する判断をAIなどの自律技術に任せるべきではなく、人間による最終判断が不可欠であるという「意味のある人間の関与」の原則を支持した文書となっている。この文書は国際社会においてLAWS規制をめぐる議論が停滞する中で、非欧米諸国による積極的な規範形成の試みとして注目された。従来、LAWSに関する国際会議は先進国主導で進められる傾向が強かったが、この共同宣言は、グローバル・サウスによる主体的な提言であった。また、この宣言により、2023年以降、カリブ共同体 (CARICOM)、西アフリカ諸国経済共同体 (ECOWAS) など他の地域機構も相次いでLAWS規制支持の声明を発表し、地域主導による多国間規範形成の動きが拡大する契機ともなった。この会議には、国連や赤十字国際委員会 (ICRC)、NGO団体 (特にStop Killer Robotsキャンペーン) なども参加し、政府・国際機関・市民社会の三者が協調して議論を進めた。

広島AIプロセス

2023年5月に閉幕したG7広島サミットの首脳宣言によって創設が盛り込まれた枠組み。AIの開発から利用に至るまでの責任とガバナンスを明確にし、国際的なルールづくりを推進している。

2023年9月には、7月～8月にOECDが起草したレポートや、生成AI等を含む高度なAIシステムの開発に関して議論すべく閣僚級会合が開催され、透明性、偽情報、知的財産権、プライバシーと個人情報保護などが優先課題であることが確認された。

その後、10月30日に「広島AIプロセスに関するG7首脳声明」が発表され、高度なAIシステムの開発者を対象とした国際指針と行動規範が公表された。

さらに同年12月には、AIに関するプロジェクトベースの協力を含む広島AIプロセス包括的政策枠組みや広島AIプロセスを前進させるための作業計画が発表されている。

ウィーン会議

2024年4月、オーストリアのウィーンで開催された国際会議「人類の岐路：自律兵器システムと規制の課題 (Humanity at the Crossroads: Autonomous Weapons Systems and the Challenge of Regulation)」には、143か国以上の政府代表や国連機関、市民団体など約1,000人が参加し、AIによる致死性自律兵器 (LAWS) の規制について議論が行われた。この会議は、国連の枠組み (CCW) 内での合意形成が停滞している状況を受け、より具体的かつ迅速な規制の枠組みを目指すために開催された。議論では、AIが人間の関与なしに人の生死を判断することへの懸念が共有され、「意味ある人間の関与 (Meaningful Human Control)」の原則を確保すべきだという意見が多くのか国から示された。また、AIの誤作動や偏ったデータに基づく判断のリスクも強調され、倫理・法・技術の観点から幅広い問題が取り上げられた。会議の成果として「ウィーン要約 (Vienna Summary)」が採択され、法的拘束力のある国際ルールが必要と、その交渉開始を求める政治的合意が形成された。

AI並びに人権、民主主義及び法の支配に関する欧州評議会枠組み条約

欧州評議会により、2024年5月17日採択。AIに関する初の国際的かつ法的拘束力を持つ条約。この条約はAI技術の進展がもたらすリスクに対応しつつ、人権、民主主義、法の支配といった基本的価値観を守ることを目的としている。

AI安全サミット

2024年5月21日から22日にかけて、韓国とイギリスの共催によってソウルで開催された。AIの安全性、革新性、包摂性を国際的に推進することを目的とした。ソウル宣言では、人間中心のAI開発の推進、国際的なAIガバナンスの協力、広島AIプロセスとの連携が強調された。

未来サミット

2024年9月にニューヨークで行われた、130ヶ国以上の代表が参加したサミット。未来に関するあらゆる事象について議論し、AIの経済的活用と公共の利益、持続可能性、国際的なガバナンスの在り方について多角的な議論がおこなわれた。このサミットによって、AIの使用や開発、理解に適切な対応を促す専門家チーム「AIパネル」が設置された。サミットで採択された「未来のための協定」ではAIを含む新興技術の誤用リスクや倫理的課題に対応するためのグローバル・ガバナンスの必要性が明記されている。とくにAIの軍事利用に関する懸念が強調されている。後日行われた「未来のための協定」採択時の国連総会ではイラン、ベラルーシ、北朝鮮、ニカラグア、スーダン、シリアの6カ国が反対票を投じた。

国際人工知能 (AI) 安全研究所ネットワーク

AIの安全性に関するグローバルな取り組みを強化することを目的としており、米国が初代議長を務めた。初期メンバーとして、オーストラリア、カナダ、EU、フランス、日本、ケニア、韓国、シンガポール、英国が参加した。AI関連の議題に関し議論が行われ、「AI安全性に関する国際協力の強化」、「合成コンテンツが及ぼすリスクに対する資金提供と研究」、「国際的なAIテストの実行結

果」、「高度なAIシステムのリスク評価に関する共同声明」、「AISI主導のタスクフォースの新設」の5点が発表された。



第2項 各国の法規制、ガイドライン等の整備動向

欧州連合(EU)

AI法

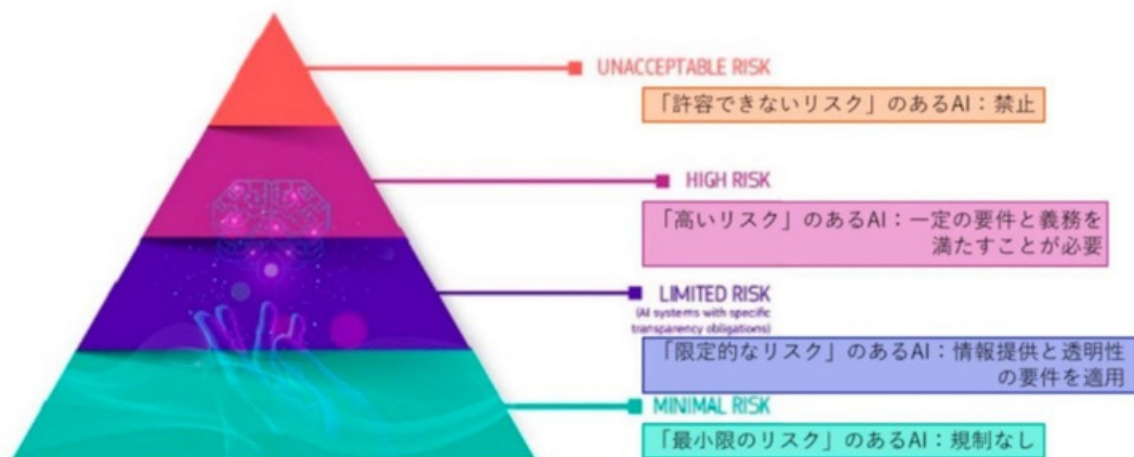
2024年3月に正式に成立したこのEU法は、EU加盟国27カ国を対象としている。

この法律は、AIの開発や利用によって生じるリスクや、倫理的な問題(たとえば偏見や差別の助長など)、さらに人間に対する被害(たとえば安全性の欠如やプライバシーの侵害)を防ぐことを目的としている。その一方で、AI技術の発展やAIによって実現される医療、教育などの社会全体の革新を妨げないように配慮されており、AI使用の安全性と技術の両立を目指している。主に人権、基本的自由の保護、AIシステムの信頼性確保、法の整合性と法の支配の維持、AI市場の統一と産業競争力の強化を目的に定められた。それによって、安全性の向上や透明性の確保、差別の防止、法的安定性が望める。

AI法は、リスクに応じて規制内容を変える「リスクベースアプローチ」という方針に基づいている。規制対象を、

- ①許容できないリスク
- ②高いリスク
- ③限定的なリスク
- ④最小限のリスク

という4段階のリスクレベルのAIアプリケーション及びシステムに分類し、それぞれに対して異なる規制を課すこととしており、上記の規制に違反した事業者には、最も重い違反の場合、最高で3,500万ユーロ(約56億円)の罰金、あるいは年間売上高の7%の制裁金が科される可能性がある。



(出典) European Commission (2024)^{*13}を基に作成

↑資料 AI法におけるリスクベースアプローチ

アメリカ合衆国

「安全・安心・信頼できるAIの開発と利用に関する大統領令」

ビッグテック企業を多く保有する米国は、自国の企業保護に力を入れ、政府による規制よりも民間での自主的な対応を優先し、企業の取り組みに任せつつ、必要に応じて政府が規制をかけるという立場をとってきた。

民間側の取り組みとして、AI開発で先行する7社（Google、Meta PlatformsやOpenAI等）がAIの安全な開発のため、新しいAIを公開する前に問題がないか検証したり、AIが作った画像や文章を人がAIが作ったものだと見分けられるようにするなどの自主的な取り組みを約束したこと、さらにその後、新たな8社（IBM、Adobe、NVIDIA等）がそれに合意したことを米国政府が発表した。各社は、自主的なコミットメントとして、安全性、セキュリティ、信頼性の3つの観点から原則を掲げている。ホワイトハウスは、強制力のある規制が導入されるまで、各社が上記の取り組みを続けるとしていた。

しかしその3ヶ月後、バイデン大統領は、「安全・安心・信頼できるAIの開発と利用に関する大統領令」を発表した。対象とするAIの問題については、従来の倫理的観点から、安全保障問題に範囲を拡充しており、対象となる事業者はビッグテック企業に限らず、バイオテクノロジー企業等、国家の安全保障や経済に影響を及ぼす可能性のあるサービスや製品を取り扱う企業も含まれる。その主な構成要素は大きく分けて3つある。

①AIに関する新たな安全性評価

NIST(国立標準技術研究所)が策定したレッドチームテスト基準に基づき、AIモデルを公開する前に厳格な安全性評価を実施する。

②公平性と公民権に関するガイダンス

司法省と連邦民権局が、AIアルゴリズムが人種・性別・宗教・障害などによって差別的な判断を下さないかを技術的に検証し、定期的に差別リスクを評価する。また、福祉給付や住宅支援などの公的制度でAIを使う際には、利用者に事前に通知する。

③AIが労働市場に与える影響に関する調査等の義務づけ

労働省にAI導入が職種構成や賃金に与える影響を調査・分析させ、その結果を報告させる。



第5章 論点解説

本章では、会議において大使のみなさんに話し合っていたきたい論点について解説している。今会議においては、以下の2つの論点からAIの規制や対応について、議論していただきたい。また、アウトオブアジェンダの内容を話題に上げるのは避け、記載された事柄を中心として話し合うようにしていただきたい。



論点1「AIに対する規制」

AIの規制に関して、各国での法的な対応は進みつつあるが、国際的なルール形成、規制は始まったばかりである。そこで大使の皆さんには、AIを国際社会で正しく活用するためには、どのような規制が必要であるのかを考えていただきたい。



生成AI

先述の通り、AIは民主主義にとってとても重要なツールであり、様々な分野で活用されている。一方で誤作動やデータが偏る懸念は尽きない。人命に関わる自動運転や医療分野などは、特に厳格な基準が必要である。

生成AIは、言語や画像を自動生成する能力を持つ一方で、フェイクニュース、偽造映像（ディープフェイク）、著作権侵害、差別的内容の拡散といった一面も持っている。これらの一面は国際的な摩擦や緊張を引き起こす原因になりうるだろう。

生成AIの国際社会における適正化、誤った使い方がされないようにするためには、どのような政策が考えられるだろうか。



軍用AI(LAWS)

特に注目されているのが自律的に行動する兵器システムだ。一般的に兵器システムの運用は、「認識」、「意思決定」、「攻撃」の三つの段階に分けることができる。現在AIは、「認識」や「攻撃」の段階での軍事利用がすでに行われている。例えば、画像認識技術を用いて標的を自動的に特定する機能は多くの兵器システムに組み込まれている。「意思決定」の段階については、依然として人間による介入が不可欠な状況である。

2024年時点で、先述の通り、国際人道法では無差別に殺傷する兵器全般を禁止しており、自律型致死兵器システム(LAWS)も無差別に攻撃をするリスクがある場合はこれに抵触する。ただ、技術的に目標を指定できる場合には該当せず、LAWSを合法的に使用することは可能となっている。

2018年以来、アントニオ・グテーレス国連事務総長は、「致命的な自律兵器システムは政治的に受け入れられず、道徳的に忌まわしいものである」として、国際法のもとでの禁止を求めてきた。そして、2026年までに、人間の制御や監督なしに機能し、国際人道法に準拠して使用できない致命的な自律型武器システムを禁止し、他のすべてのタイプの自律型武器システムを規制する法的拘束力のある文書を締結するよう、各国に勧告した。

LAWSを規制するための国際法については2025年現在、検討が進められている最中だ。

LAWSは国際社会においてどのような扱いをするべきものであるのか、完全禁止するべきなのか、それとも使用条件などを設けた規制をすべきなのかなど、各国の国益をぶつけ合う激しい議論を期待している。



論点1のまとめ

- ①生成AI 生成AIの国際社会における適正化、誤った使い方がされないようにするためには、どのような政策、規制が考えられるだろうか。
- ②軍事用AI LAWSは国際社会においてどのような扱いをするべきものであるのか、完全禁止するべきなのか、それとも使用条件などを設けた規制をすべきなのか



論点2「AIの誤作動に対する責任の所在」

AIの誤作動によって引き起こされる問題は、生成AIや軍事用AIを問わず、すでに現実的な懸念として浮上している。たとえば、生成AIが虚偽情報を拡散し、国際的な誤解や緊張を引き起こすケースや、軍事用AIによるターゲティングミス(標的の誤認)によって民間人に被害が及ぶケースなどが想定される。

AIによる意思決定の際には、誰が最終的な責任を負うのかという点を明確にすることが欠かせない。AIによる事故や失敗が発生した際、どこに求めるべきかが明確ではない現状がある。そして、このあいまいさがさらなる技術の利用拡大を阻害する要因の一つとなっている。

ここでAIの責任帰属について、一般的な考え方をいくつか例示しようと思う。

○開発者(AIの設計者、開発企業)

AIの行動はその設計に依存しており、誤作動は設計ミスや想定不足の結果であるとする考え方。

○利用者(AIを使用した人、組織)

利用者の監督責任であるとする考え方。

○提供者(AIを提供、販売した企業、プラットフォーム運営者)

生成AIのようにAIが不特定多数に利用される場合、規制や制限を行わなかったプラットフォーム側に責任を問う考え方。

○共同責任(開発者・提供者・利用者など複数の関係者が分担)

AIの行動は設計・提供・使用の複合的な過程の結果であるとする考え方。法制度を設計する上で、最も現実的であり主流になりつつある。

○国家

軍の行動は国家の行動であり、兵器や権力行使は国家が責任を持つべきであるとする考え方。国連国際法委員会(ILC)が定めた「国家責任に関する条文(Articles on State Responsibility)」では
国家の行為によって他国に損害が生じた場合、その行為が国家に帰属する限り国家が責任を負うとされている。

○AI自身

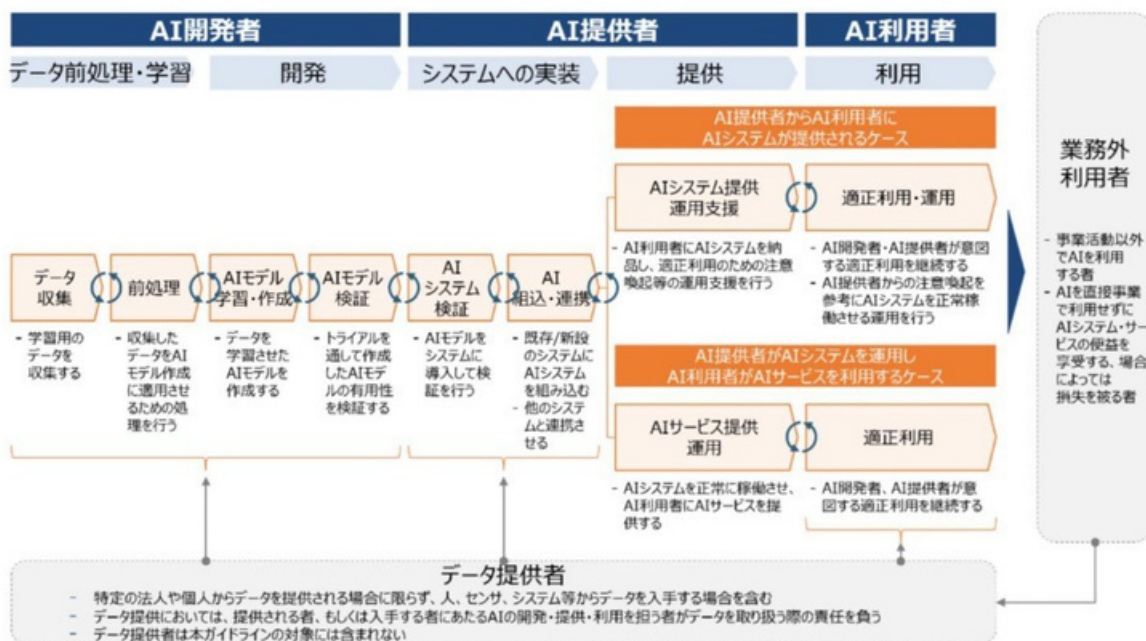
AIに法的人格や自己資産を持たせ、賠償などの法的行為を可能にしようとする考え方。

AI開発者は、AIモデルを直接的に設計し変更を加えることができるため、AIシステム・サービス全体においてもAIの出力に与える影響力が高い。また、イノベーションを牽引することが社会から期待され、社会全体に与える影響が非常に大きい。このため、自身の開発するAIが提供・利用された際にどのような影響を与えるか、事前に可能な限り検討し、対応策を講じておくことが重要となる。

AI提供者は、AI開発者が開発するAIシステムに付加価値を加えてAIシステム・サービスをAI利用者に提供する役割を担う。AIを社会に普及・発展させるとともに、社会経済の成長にも大きく寄与する一方で、社会に与える影響の大きさゆえに、AI提供者は、AIの適正な利用を前提としたAI

システム・サービスの提供を実現することが重要となる。そのため、AIシステム・サービスに組み込むAIが当該システム・サービスに相応しいものか留意することに加え、ビジネス戦略又は社会環境の変化によってAIに対する期待値が変わることも考慮して、適切な変更管理、構成管理及びサービスの維持を行うことが重要である。

AI利用者は、AI 提供者から安全安心で信頼できるAIシステム・サービスの提供を受け、AI提供者が意図した範囲内で継続的に適正利用及び必要に応じてAIシステムの運用を行うことが重要である。これにより業務効率化、生産性・創造性の向上等AIによるイノベーションの最大の恩恵を受けることが可能となる。また、人間の判断を介在させることにより、人間の尊厳及び自律を守りながら予期せぬ事故を防ぐことも可能となる。(AI 事業者ガイドライン - 経済産業省引用)



↑資料 一般的なAI活用の流れにおける主体の対応

AIの誤作動に関する責任の所在は、未だ国際的な基準が存在しない。そのため大使の皆さんには、AIの誤作動による被害に対する責任の所在について、どのような国際的基準を設ければ良いか話し合っていたきたい。上記のように、開発者、提供者、利用者はそれぞれの形でAIに関わり、影響を与えている。誰がコントロールしていたのか、その誤作動は予測できるものだったのか、AIは自律的に意思決定を行なったのか、実際に責任を負うことが出来るのは誰か、などAIの責任問題を考える上での観点はいくつも存在する。また、実際に国際的な被害を生み出してしまった場合、どこで制裁に関する議論を行えば良いだろうか。AI被害に特化した国際的なAI責任調整機関を設立すべきなのだろうか？それとも、既存の機関で良いのだろうか？

AIの誤作動が、人命や社会に直接的な影響を及ぼす可能性が否定できない現代において、このような議論は避けては通れないものとなっている。AIに関わる全ての立場から責任問題を捉え、様々な観点から意見を出していただきたい



論点2のまとめ

- ①AIの誤作動による被害に対する責任の所在について、どのような国際的基準を設ければ良いか。
- ②AIの使用によって実際に国際的な被害を生み出してしまった場合、どの機関で制裁に関する議論を行えば良いだろうか。



アウトオブアジェンダ

アウトオブアジェンダとは会議における議論で基本的に避けるべき項目である。各国の事情で少し触れてしまう場合もあるだろうが、決議文書に載せることは不可能であることに留意していただきたい。



AI導入に対する議論

AI技術は今後、発展途上国を含むすべての国において不可避な影響を及ぼすものであり、その導入や活用は将来的に避けられない課題である。ゆえに、本会議では経済支援等の短期的・補助的な議題ではなく、AIに対する規制や責任の所在といった根本的かつ長期的視点からの制度的課題に焦点を当てて議論を行うものとする。



AIの人権問題に関する議論

論点2では、AIの利用によって生じた問題に対する責任の所在について議論していただく。特に、「AI自身に責任を取らせる」という選択肢についても議論の対象とするが、これはAIに人権や意識、人格を認めるという倫理的・哲学的な立場に基づくものではない。

ここでの「責任を取らせる」という概念は、AIを法人格に類する制度的な主体として捉え、開発者・運用者・企業など既存の人間主体との責任配分をより明確化するための法的枠組みとして議論するものとする。

よって、大使のみなさんにはAIの権利や感情といった人権に関わる議論は含めず、責任の制度的設計に焦点を当てて議論を進めていただきたい。



過度に専門的な議論



参考文献

この議題概説書を作成するにあたって参考にした文献やウェブサイトを記載している。必要に応じて閲覧し、リサーチの手がかりとして活用していただきたい。

国連文書

https://www.unic.or.jp/info/un/charter/text_japanese/

公式ページ(UNODA)

[Convention on Certain Conventional Weapons \(CCW\) – UN Office for Disarmament Affairs](#)

ウェブサイト

患者と臨床医にとっての最新AI技術の利点

[The Benefits of the Latest AI Technologies for Patients and Clinicians](#)

Benefits and Risks of AI in Health Care: Narrative Review - PMC

<https://pmc.ncbi.nlm.nih.gov/articles/PMC11612599/>

AIロボットは日本の高齢化社会の介護の鍵を握っているかもしれない

[Work-Life Balance - OECD Better Life Index](#)

AI robots may hold key to nursing Japan's ageing population

[AI robots may hold key to nursing Japan's ageing population | Reuters](#)

Work-Life Balance - OECD Better Life Index

[Work-Life Balance - OECD Better Life Index](#)

Japan firms face serious labour crunch from aging population ...

[Japan firms face serious labour crunch from aging population ...](#)

[社説]AIが殺傷決める兵器に世界で歯止めを - 日本経済新聞

<https://www.nikkei.com/article/DGXZQODK2967D0Z20C24A8000000/>

民間企業の取り組み

ホワイトハウス公式サイト

https://bidenwhitehouse.archives.gov/briefing-room/statements-releases/2023/07/21/fact-sheet-biden-harris-administration-secures-voluntary-commitments-from-leading-artificial-intelligence-companies-to-manage-the-risks-posed-by-ai/?utm_source=chatgpt.com

アメリカ政府のAI安全対策は本当に前進なのか——それともまたしても歯の抜けた約束なのか？

https://www.techpolicy.press/the-us-governments-ai-safety-gambit-a-step-forward-or-just-another-voluntary-commitment/?utm_source=chatgpt.com

AI's mysterious 'black box' problem, explained

[AI's mysterious 'black box' problem, explained](#)

無人兵器システムにおける自律性(Autonomy)の利点と欠点

[無人兵器システムにおける自律性\(Autonomy\)の利点と欠点](#)

Lethal Autonomous Weapon Systems (LAWS)-UNODA

<https://disarmament.unoda.org/the-convention-on-certain-conventional-weapons/background-on-laws-in-the-ccw/>

恐怖のAI兵器「LAWS(自律型致死兵器システム)」とは？使用規制が進まない複雑な事情

<https://www.sbbbit.jp/article/cont1/152252>

国際人道法とは - 赤十字国際委員会

<https://jp.icrc.org/information/what-is-international-humanitarian-law/>

国際人道法の重要要素

<https://jp.icrc.org/activity/israel-and-the-situation-in-gaza-from-the-perspective-of-international-humanitarian-law/>

「第3の軍事革命」AI兵器 軍事利用加速で“戦争のカタチ”変わる？

[「第3の軍事革命」AI兵器 軍事利用加速で“戦争のカタチ”変わる？](#)

WEB 特集「AI兵器」の衝撃 “機械は犠牲を理解できず”暗い未来の不安

[WEB 特集「AI兵器」の衝撃 “機械は犠牲を理解できず”暗い未来の不安
の不安](#)

European Commission, "AI Act"

<https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>

AIに関する各国の対応-総務省

<https://www.soumu.go.jp/johotsusintokei/whitepaper/ja/r06/pdf/n1420000.pdf>

AI時代の法律：知財、契約、責任問題を徹底解説

<https://riter.bring-flower.com/articles/ai-legal-intellectual-property-contracts-liability/>

AI事業者ガイドライン-経済産業省

https://www.meti.go.jp/shingikai/mono_info_service/ai_shakai_jisso/pdf/20240419_1.pdf

AIと自律性の責任ある軍事利用に関する政治宣言「軍事領域における責任あるAI利用(REAIM)」サミットより

<https://www.mofa.go.jp/mofaj/files/100580933.pdf>

AIドローンの誤作動が戦争の悲劇を招く——元グーグルの開発担当が警鐘

https://forbesjapan.com/articles/detail/29731?read_more=1

AIの軍事利用とは？用いるメリットや問題点も詳しく解説

<https://ai-kenkyujo.com/news/ai-gunziriyoku/#:~:text=%E5%88%B6%E5%BE%A1%E3%81%8C%E3%81%A7%E3%81%8D%E3%81%AA%E3%81%84%E6%81%90%E3%82%8C%E3>

[%81%8C.%E6%80%A7%E3%81%8C%E3%81%82%E3%82%8B%E3%81%AE%E3%81%A7%E3%81%99%E3%80%82](#)

EZproxy Tools - The Open University
[EZproxy Tools - The Open University](#)

AWS Diplomacy Report, Vol. 1, No. 1

[AWS Diplomacy Report, Vol. 1, No. 1](#)

オーストリア外務省:2024年ウィーン会議 公式ページ
[2024 Vienna Conference on Autonomous Weapons Systems – Austrian Federal Ministry](#)

Stop Killer Robots:ウィーン会議は新たな国際法への支持を確認

[Vienna Conference Affirms Commitment to New International Law – Stop Killer Robots](#)

人類の分岐点としてのウィーン会議

[Vienna Conference – Article 36](#)

ヒューマン・ライツ・ウォッチ:LAWS規制条約への圧倒的支持

[Resounding Support for Killer Robots Treaty – Human Rights Watch](#)

Stop Killer Robots:中南米33カ国がベレン共同コミュニケを発表

[Significant Act of Political Leadership – Stop Killer Robots](#)

HRW:中南米諸国がLAWSに反対

[Latin America and Caribbean Nations Rally Against Autonomous Weapons – Human Rights Watch](#)

Automated Research:法的拘束力ある規制交渉を求めるベレン共同声明
[Belen Communique Calls for Urgent Negotiation – Automated Research](#)

国連軍縮局:特定通常兵器使用禁止制限条約(CCW)とLAWS

[UN Office for Disarmament Affairs – LAWS page](#)

武器の規制と国際人道法の観点からLAWSを批判的に分析。

Article 36 – Regulation of weapons

LAWS規制を訴える国際NGO。各会議の解説や政策提言あり。

[Stop Killer Robots – Global campaign to ban autonomous weapons](#)

「国連と赤十字が自律型兵器の規制を共同で呼びかけ(UN News, 2023年10月)」

<https://news.un.org/en/story/2023/10/1141922>

「殺人口ロボット: 国連新報告書、2026年までの条約必要と訴える(HRW, 2024年8月)」

<https://www.hrw.org/news/2024/08/26/killer-robots-new-un-report-urges-treaty-2026>

「グテーレス国連事務総長、AI兵器に関する国際規制を要請(NHK, 2024年8月)」

https://www3.nhk.or.jp/nhkworld/en/news/20240827_06/

「国連軍縮会議、AI兵器議論の拡大を検討(Arms Control Today, 2024年12月)」

<https://www.armscontrol.org/act/2024-12/news/un-moves-expand-autonomous-weapons-discussions>

「国連第一委員会、LAWSに関する新決議を可決(UN Press, 2023年11月)」

<https://press.un.org/en/2023/gadis3731.doc.htm>

「国連軍縮機関、LAWSへのグローバル対応を呼びかけ(Computer Weekly, 2023年10月)」

<https://www.computerweekly.com/news/366559035/UN-disarmament-body-calls-for-global-action-on-autonomous-weapons>

「2025年5月 国連で“殺人口ロボット”規制が遅れていることへの懸念(Reuters)」

<https://www.reuters.com/sustainability/society-equity/nations-meet-un-killer-robot-talks-regulation-lags-2025-05-12/>

「国連: 殺人口ロボット禁止条約の交渉を開始すべき(HRW, 2025年5月)」

<https://www.hrw.org/news/2025/05/21/un-start-talks-treaty-ban-killer-robots>

「国連: 殺人口ロボット規制交渉への支持広がる(HRW, 2024年12月)」

<https://www.hrw.org/news/2024/12/05/killer-robots-un-vote-should-spur-treaty-negotiations>

「殺人口ロボット禁止の国連投票、条約交渉の後押しに(HRW, 2024年1月)」

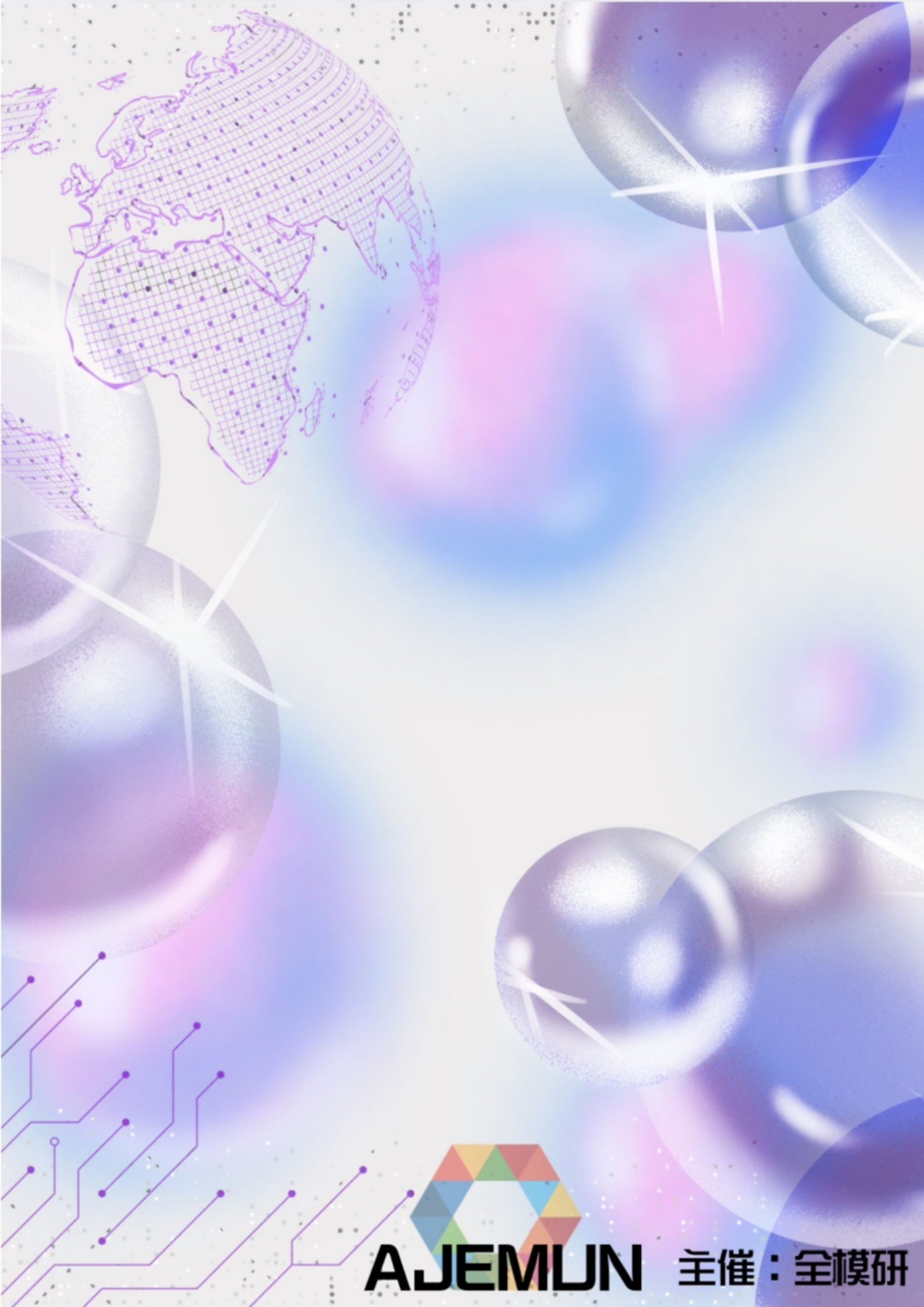
<https://www.hrw.org/news/2024/01/03/killer-robots-un-vote-should-spur-action-treaty>

「特定通常兵器使用禁止制限条約(CCW)の公式ページ」

<https://www.un.org/disarmament/the-convention-on-certain-conventional-weapons/>

「2021年 GGE報告書(意味ある人間の関与などを盛り込んだ画期的文書)」

https://documents.unoda.org/wp-content/uploads/2022/01/CCW_GGE.1_2021_CRP.2-Advanced-unedited-version.pdf



AJEMUN

主催：全模研