

All Japan Educational Model United Nations



United Nations
General Assembly
1st Committee (DISEC)

EIGHTIETH UNITED NATIONS GENERAL ASSEMBLY FIRST COMMITTEE

A/80/1/DR.1

Agenda item: AI と軍事 (AI and the military)

2025 年 8 月 5 日

Sponsor: Argentina, Austria, Bangladesh, Brazil, Cambodia, Canada, Chile, Costa Rica, Egypt, Estonia, Ethiopia, Germany, Iran, Italy, Japan, Jordan, Kenya, Latvia, Lebanon, Lithuania, Mali, Mexico, New Zealand, Norway, Oman, Pakistan, Panama, Philippines, Portugal, Serbia, South Africa, South Sudan, Sudan, Tajikistan, Ukraine, United Kingdom

第 80 回国連総会第一委員会は、

人工知能(AI: Artificial Intelligence)を使用した大量破壊兵器を認めないことを強調し、

1948 年に国連総会で採択された「世界人権宣言」を想起し、

1949 年の「戦争犠牲者の保護のためのジュネーブ諸条約 (Geneva Conventions for the Protection of War Victims)」と 1977 年に締結された二つの追加議定書 (1977 Additional Protocols) に、武力紛争の際に適用する原則や規則を網羅している「国際人道法」に好意を示し、

国連総会によって設置された「政府間専門家会合 (GGE)」の取り組みに賛同し、

自律型致死兵器システム (以下、LAWS: Lethal Autonomous Weapon Systems) は AI によって自律的に攻撃できる、非人道的な兵器であることを遺憾に思いながら言及し、

軍事的な AI の使用において意味ある人間の関与が必要であることを断言し、

LAWS には、必ず人間の意志で使用を中断できるプロセスが無くてはならないことを強調し、

AI の正しい使用を促進させるためには、国内、国外共に取り締まりや協力体制を強化する必要があることを認識し、

AI にかかわる人々は責任感を持って行動しなければならないということを再確認し、

AI などの先進技術は、生活を豊かにするために開発されたものであるということを再確認し、

健全な AI の発展は国の利益になることを認識し、

人の生命や基本的人権に対して脅威をもたらす AI システムの使用は望ましくないということを強調し、

AI が持つ持続可能な開発への潜在的な貢献を認識するとともに、その誤用によるリスクを認め、

生成 AI のシステムのサイクル全体における透明性、責任性、人権尊重、安全保障の重要性を強調し、

将来、AI 技術がより拡充した際に既存の機関での AI 問題解決が難航する可能性を認識し、

今後基準の改訂のための議論を続ける必要性を認識し、

各国および国際機関による AI 規制枠組み構築への継続的な取り組みを是認し、

過去の国際会議や条約に強制力がなく、いまだに全ての国が準拠する軍事用 AI に関する国際的な基準や議論の場所が存在していないことを再認識し、

軽率な LAWS の使用の結果として、不測の民間人被害、武力行使の促進、人道的危機の拡大を招く可能性があることに深い憂慮を表明し、

軍事 AI の判断ミスで人的被害が出るリスクがあることを認識し、

今まで新兵器が開発された時に有効な国際基準の制定が間に合わなかったことを起因に、多大な被害を出してしまったことを反省し、

各国内の政治は、各国の中核となる行為であり、その国民のために独自に行っていくものであると信じ、

いかなる場合でも生成 AI によって人権侵害が行われてはならないことを再確認し、

ディープフェイクが国際社会にリスクを与えることを確認し、

1. 各国に対し、AI が誤差動を起こした際に被害が生じるリスクをベースとした下記のような基準で AI を評価し、規制するように要請する：
 - a. 許容できないリスクを持つ AI に対してはその利用を禁止、
 - b. 司法やインフラに使用されているような高いリスクを持つ AI に対しては以下のように規制：
 - i. 人間による監視、
 - ii. サイバーセキュリティの確保、
 - c. ディープフェイクなどの、限定的なリスクを持つ AI に対しては作成元の透明性を確保する義務を課すこと、
 - d. 上記以外、最小限のリスクを持つ AI に対しては特に義務を設けず、各国が主体性を持って規制する；
2. 加盟国に対し、生成 AI で生成したものには、それがわかるようにラベルなど、証明になるものを表示するよう要請する；
3. 加盟国に対し、完全自律型兵器、及び、最終判断を AI に任せ、人間の意思決定なしに殺傷力をもって交戦することのできる兵器システムの開発、配備および使用を全面的に禁止することを強く要請する；
4. 加盟国に対し、軍事分野における AI システムに「意味のある人間の関与（meaningful human control）」を適用し、以下の事項に準拠することを要請する：
 - a. 開発時、評価基準に則った製造を行い、誤作動がないか検証を行うこと、
 - b. 緊急停止ボタンのような人間の意思で使用を中断することのできるプロセスを内蔵すること、
 - c. 誤作動やハッキングを防ぐために、常にサイバーセキュリティを確保すること、
 - d. 使用目的を防衛のみとし、民間人に危害を加えないこと、
 - e. あくまで国の管理下に保持し、非国家団体に渡らないようにすること、
 - f. 最終的な判断は人間が行うこと、

- g. 兵器が運用されている間、人間が常時情報収集を続けること,
 - h. 人間による誤作動を防ぐためのフルプルーフ機能を搭載すること,
 - i. 誤作動を起こした際の責任の所在を明らかにすること;
5. 国際連合に対し、AI の取り締まり・諸問題への対処等の以下の業務を行うための新機関として、United Nations Agency for the Joint Enforcement and Monitoring of the Use of Neural AI (UNAJEMUN)を設立することを要請する:
- a. AI の誤作動を起こした際の責任追及のために、AI の誤作動による被害の責任帰属に関し、あらかじめ決められている責任を鑑み、その誤作動の原因に応じて、複合的に責任を追及するという前提の上で、責任者への制裁を行うための裁判を行うこと,
 - b. LAWS の誤作動を防ぐためのガイドラインや、開発時・使用時における評価基準の作成,
 - c. a における裁判で用いる規範の作成,
 - d. 各国の AI の使用状況等に関する報告書の収集,
 - e. AI の適正利用に関する定期的な会議の実施,
 - f. AI の誤作動による重大な被害発生時の調査の実施,
 - g. 完全自律型兵器が継続的に禁止されていることを確認するための各国への監視,
 - h. 特定の AI について、以下の点が担保されているかの判断:
 - i. 公平性,
 - ii. 透明性,
 - iii. 説明責任,
 - iv. プライバシーの尊重;
 - i. ディープフェイクを見分ける技術の提供,
 - j. AI の誤作動による被害が出た場合のその AI の規制の内容の協議;
6. 各国政府に対し、核兵器、生物兵器、化学兵器を含む大量破壊兵器に AI 兵器を接続しないことを強く促す;
7. 各国政府に対し、新機関に、自国の LAWS 開発と規制の現状を一年に一回報告することを要請する;
8. 各国政府に対し、半自律型兵器の使用前と使用後に CCW に以下のことを報告することを促す:
- a. 使用する目的,
 - b. 使用する手段,
 - c. 使用した結果;
9. 主文14番について、各国が長期的な使用を行う場合は定期的な報告を行うことを促す;
10. 国際社会に対し、国内の選挙を各国にとって適切な形で行っていくために以下のことを促す:
- a. 生成 AI を使用して国の選挙に影響を与えようとする行為の禁止,

- b. 他国の生成 AI によって作成された情報によってなんらかの影響を受けたことが疑われる、または特定された場合に、それを他国に対して主張することの容認；
- 11. 各国政府に対し、生成 AI によってプライバシーと個人情報が漏洩しないようにするよう促し、このことの確認を新機関で行うことを要請する；
- 12. 各国政府に対し、ディープフェイクで生成されたものが判別できるようにすることを促す；
- 13. 各国政府に対し、特定の AI が、以下の点が担保されていないと判断された場合、規制対象とするよう要請する：
 - a. 公平性,
 - b. 透明性,
 - c. 説明責任,
 - d. プライバシーの尊重；
- 14. 全加盟国に対し、生成 AI が出力する結果に情報の偏りがないよう以下の内容の実施を推奨することを奨励する：
 - a. 社会的マイノリティについて AI の学習量と質を向上させる,
 - b. 少数言語での AI の学習量を増加させ、出力の精度を向上させる；
- 15. 各国政府に対し、新しい AI を開発する際、開発者はその AI のあらゆる可能性や危険性をまとめた文書を決まった機関に提出することを強く要請する；
- 16. 各国政府に対し、AI の誤作動によって被害が出た場合、責任の所在について：
 - a. 開発者が提出した文書の範囲内であった場合は監督不足として、利用した企業・軍に責任を追及することに賛成し,
 - b. 文書に記載がなかった場合は実験不足や想定不足であったとして、開発者に責任を追及することに賛成する.