

N P まとめ

A議場

A01



Argentina

論点1

- **生成AIに関する緩やかな規制**

誤用をどう防止するか？

→ **自主規制や倫理ガイドラインの強化、業界の責任ある開発**で対応

- **LAWS全面規制**

「意味のある人間の関与」、つまり**人間による最終判断を重視**

論点2

- **LAWSは軍事用AIにおいてAIの最終判断により殺傷能力を行使できるもの**だと定義

- AIの使用により被害が生まれた場合、**制裁は国連安全保障理事会において議論**を行う

- AIの誤作動による被害の責任は、関係者が分担しあう**共同責任の考え方を世界共通で採用する**



A議場 AUSTRALIA



ごきげんよう！オーストラリア大使です

まず生成AIに関して、プライバシー保護を進めるため、ガイドラインを作りたいです。

次に軍用AIに関して、**国際人道法**と**倫理**を重要視しながら使用したいです。

2日間、よろしくお願いします！



アナザーストーリー R2-D2の謀反

遥か彼方の銀河系――

英雄ルーク・スカイウォーカーが焼死体で発見された
専門家によれば、犯人はかつて忠実な相棒だった完全自立型AI、
「R2-D2」であった。

検察は彼を殺人の容疑で起訴した。

彼の主張はこうであった「故意ではない」と。

裁判では彼の主張が認められ罪名は「傷害致死」で有罪。罪は
開発元「インダストリアル・オートマン」が負うこととなった。

この星ではAIの暴走が相次いでいる。

別の事件では、半自律型AI「C-3PO」が虫と人間を誤認し、
「ドアの前にいる蜂を攻撃しますか?」と尋ねた末、宅配業者
を焼き殺した。

本件では殺人の罪で使用者であるアナキン・スカイウォーカー
が起訴され、有罪となった。

次に暴走するAIは遥か彼方ではなく、もっと近くで起こるかも
しれない――

(※上記ストーリーは、本文書の執筆者の創作です)

我々オーストリアは軍事用AIに対し以下のような政策を
掲げます。

完全自立AIの完全禁止、半自律型AIの適切な使用規制
そして半自律型AIを用いる場合、責任は全て使用者やそ
の使用

組織(国家など)が負い、適切な制裁を受けます。

C-3POが暴走する前に、この惑星を守りましょう!

Bangladesh

論点1

生成AI及び軍事用AIに関する
倫理・規制の枠組みの制定

論点2

AIによる事故や誤作動
が起きた際の対処法

AIの安全な運用のため、しっかりとした責任の所在の明確化や事故・誤作動への対処法が必要です。これらの問題の解決も考慮しつつ、議論を重ね共通方針の作成を目指しましょう

Burkina Faso

論点1

- ① リスク・ルール・技術標準ベースアプローチの3つを合わせた法的拘束力が強い規制の設置
- ② 全面禁止

論点2

- ① 責任の所在は共同責任
- ② 「国際AIガバナンス機関 (IAGO)」の創設

皆さんへ

ブルキナファソはAIが発達段階にあるため、柔軟な対応をとることが可能です。

誰一人取り残さない

-leave no one behind-

を実現するため政策が異なっても積極的に交渉していきましょう。

お気軽にお声がけください！



CAMBODIA



OUR STANCE

- 生成AIは今後我が国の国家開発において必要不可欠な存在になるので、使用は継続していきたいと思います。
- カンボジアを拠点としたサイバー攻撃やハッキング犯罪が増えている以上、社会の治安維持のために具体的な規制が必要と考えます。
- 倫理的な観点により、LAWSについては責任の所在を確定させましょう

POINT #1

AIの開発・サイバー犯罪における歴史的
背景及び人権保護上の規定



POINT #2

AI使用における責任者を事前に把握する
登録制度の設立





CANADA



<論点1>

AI軍事利用の無秩序な拡大による安全保障リスクを抑制し国際秩序を維持するため、先進国と協力して国際的な規制枠組みの整備を強く支持する

<論点2>

AI誤作動時の迅速かつ透明な情報共有と国際協力体制の確立を提案し、機密保持に配慮しつつ軍事的誤解防止と危機管理能力向上を図る





A09

Central African republic

AJEMUN2025 8.4~8.5

論点 1 :

生成AI:使用した際の表示義務

高リスク分野での使用制限

軍用AI:完全自律型兵器の禁止

→機械インタフェースを義務付け

論点2:

生成AI、軍用AIともに問題検出後、

初期診断、被害対応、分析・原因究明による

対策実行

二日間を有意義な時間にしましょう 1



Chile



【 チリが重要視している主な2つの問題 】

AI 兵器が人間の判断なしで生き物の命を奪う可能性

軍事目的ばかりに使われてしまう可能性

AI が良い方向に使われる安全な世の中に変えていきましょう。
皆さんと一緒に話し合えることを楽しみにしています！

(論 1) **LAWS**について:

意味ある人間の関与を重視

提示する定義:「人間の手が加わることのない、または自身で発展する余地がある兵器」

⇒**完全禁止**を、**議場全体**で

⇒LAWSを禁止する**国際条約**の作成の
手続きをこの会議にて開始する



(論 1) 生成**AI**について:

情報セキュリティと**主権**を大事にし
ながら、今後も発展を促す

CHINA



COSTA RICA

論点1

- 生成AI利用推進派 ↓だけど
- 生成AI利用に関してのガイドライン統一(情報モラル統一、犯罪防止ため)
- いかなる状況においても軍事用AIの利用の禁止 (LAWSも国際人道法違反として完全規制の方針)

論点2

- 明確な責任三層の分担(開発者、一般人、国家で分類し、それぞれの責任を定める)、状況に応じた責任追及
- 国際的なAI緊急対応ネットワークの創設(誤作動発生次第情報報告、原因究明、責任調査を行う)

皆が納得して実現できる決議案を作しましょう!

皆さんと実りある議論が作れることを

楽しみにしています!

A13

Egypt



☆論点1

リスクベースアプローチに則ったAIの開発における審査基準を制定する。

軍事用AIに関する**法的拘束力を持つ規制**を制定する。

☆論点2

立場によって比重の違う**共同責任**

誤作動の起きたAIの**再開発**

軍事用AIには強く反対！

発展途上国として**AI**の恩恵を平等に享受できる国家を目指します！





Eritrea

論点 1

**安全保障の不均衡是正と
技術アクセスの公平性**

論点 2

**AI軍事技術の倫理・人道的規制
と地域安定への配慮**

A15



Republic of Estonia

POINT :

論点 1 「不純物」の除去

論点 2 「監視機関」の設立

2 日間を駆使して有意義な会議に

しましょう！



AI AND MILITARY



小国と途上国の安全保障に貢献

大国によるAI乱用の被害を最小限に



フィンランド

FINLAND

(論点1)

適切な規則整備によりAI利用の安全性を保つ
生成AI

- ①各国に法的制限を設けるように促す
- ②失業防止対策
- ③国際的なAI教育の推進

軍事用AI

- ①人の関与の必要性を強調
- ②軍事用AI運用に関する試験基準の策定
- ③誤作動を防止するための対策

(論点2)

AIに関連する人々の連帯責任とする

2日間よろしくお願いします！！





FRANCE

論点①：AI兵器の使用・規制について

人間が最終判断を行う仕組みを国際ルール化すべき
AIは「補助的な道具」であり、命の判断を機械に任せてはいけない！

論点②：AIの誤作動や事故への対応

AI事故を国際的に記録し、責任を明確にする制度の
設立

AIは「補助的な道具」であり、命の判断を機械に任せてはいけない！

「私たちはAIを“恐れる”のではなく、“責任を持って
使う”道を選びます」

皆さんの知恵と経験を、この新しい国際ルールづくりに
活かしませんか？ 



Germany

Germanyは論点を軍事と民間に分けて議論し、

・民間

説明性(explicitly)をAIに含めることを奨励し、
ブラックボックスを無くす方向へ動かす。

AIの誤作動への責任はこの説明による関係者の連帯
した責任とするが、説明を果たせなかった場合は企
業が全ての責任を負う。

・軍事

LAWSを民間人に対する大きな脅威と認識し、
ナパーム弾同様民間地域での使用を禁ずる。

また、LAWSが特に紛争地帯で猛威を振るうであ
ろうことを認識し、十分な法整備が整うまで輸出
禁止モラトリアムを実施する。

良い会議にしましょう！

インド



論点1

**AIのリスクに応じた
ガイドラインを！
AIを扱う作業部会の
設置**

論点2

**ガイドライン作成
複数のステークホルダー
が責任を負うべき！**

2日間よろしくお願いします！

A22

Indonesia



**我が国では国家 AI戦略を
進めてAIの開発者として
国際平和に協力していきます**

論点1

AIを使う上での法的結束力のある規制を
設ける

論点2

AIの作用や管理方法に関する記録方法の強化

**AIの安全な使用のために
国同士で協力していきましょう！**

AIの発展は



止めてはいけない



Italy

論点①

- ・AI法のワールドスタンダード化
- ・LAWSの完全禁止

論点②

AI対応専門の新機関の設立

**AI規制には強硬姿勢で
臨みます。**

**2日間有意義な時間に
しましょう！**

論点1ーAIに対する規制ー

- ・民間人に対する被害のリスクがある都市部でのLAWS使用を禁止する。
- ・AIに関する情報共有及び議論、ノウハウ共有を行う国際的な専門の枠組みを設置し各国にAI規制法の整備を促す。

論点2ーAIの誤作動に対する責任の所在ー

- ・AIの誤作動による被害の責任帰属は責任所在が明らかな場合を除き、当該AIと開発者、使用者、運用者など複数の関係者にわたってあるとする。
- ・論点1での枠組みの下で責任追及についての議論を行う。
- ・上記枠組みのもとにAIの誤作動の原因を調査、国際社会に共有する。

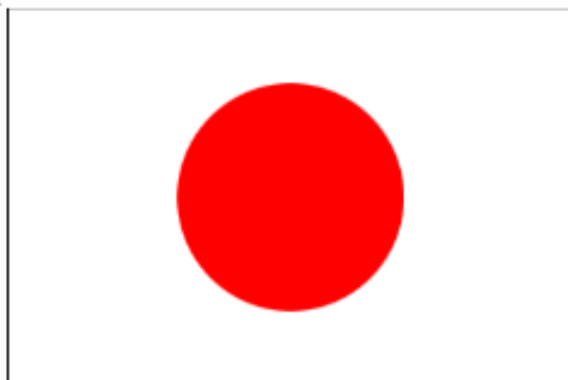
特に私たちは論点1を重視し
上記の政策を私たちは提案します
有意義な二日間にしましょう！



Japan



Here



論点1

生成AI：イノベーション促進と安全確保の両立

軍事用AI：「責任ある人間の関与」と国際人道法の遵守の重要視



論点2

生成AI：・国際ルールの制定

・技術革新と安全性のバランスを重視

・誤作動による差別・偏見・誤情報の拡散から
人権、民主主義の価値観を守る

軍事用AI：・国際人道法に基づいた責任の所在の明確化

・透明性のための制度化

AIの進歩が急速に進む今こそ、
国際的に一致した指針を作成するべきです。
世界全体で協力してよい会議にしましょう！

Jordan



論点1

- ・ヨルダンでは若者の失業率が高く、AI産業の発展が雇用創出の鍵。
- ・教育・訓練を通じたAI人材の育成が急務。

論点2

- ・軍事活用のAIは発展途上だが、監視・情報分析などで活用する可能性あり。
- ・中東の安定のため、AIは国境管理などに有効。

目標

「安全保障×経済発展」の両立を目指し、
国際協力と人材育成を通じて、平和利用を進めましょう！

Kuwait

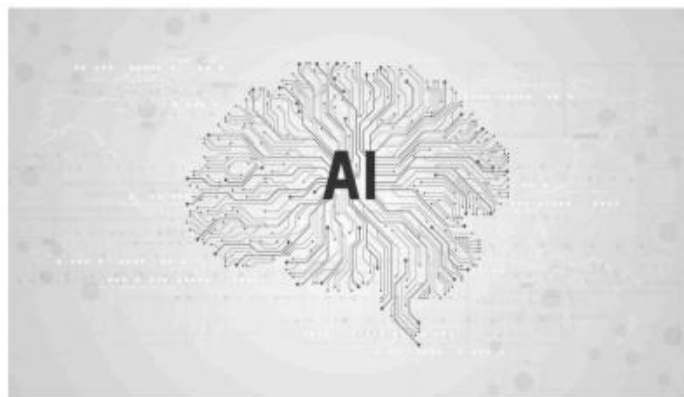


論点1

生成AIのプライバシー保護の強化

論点2

責任帰属の国際条約制定



2日間よろしくお願いします！



Latvia

A30

論点1
AIに
対する
規制

論点2
AIの誤作動
による
責任の所在

多国間協調性と透明性のある会議
にしていきたいです。

Lebanon



論点①

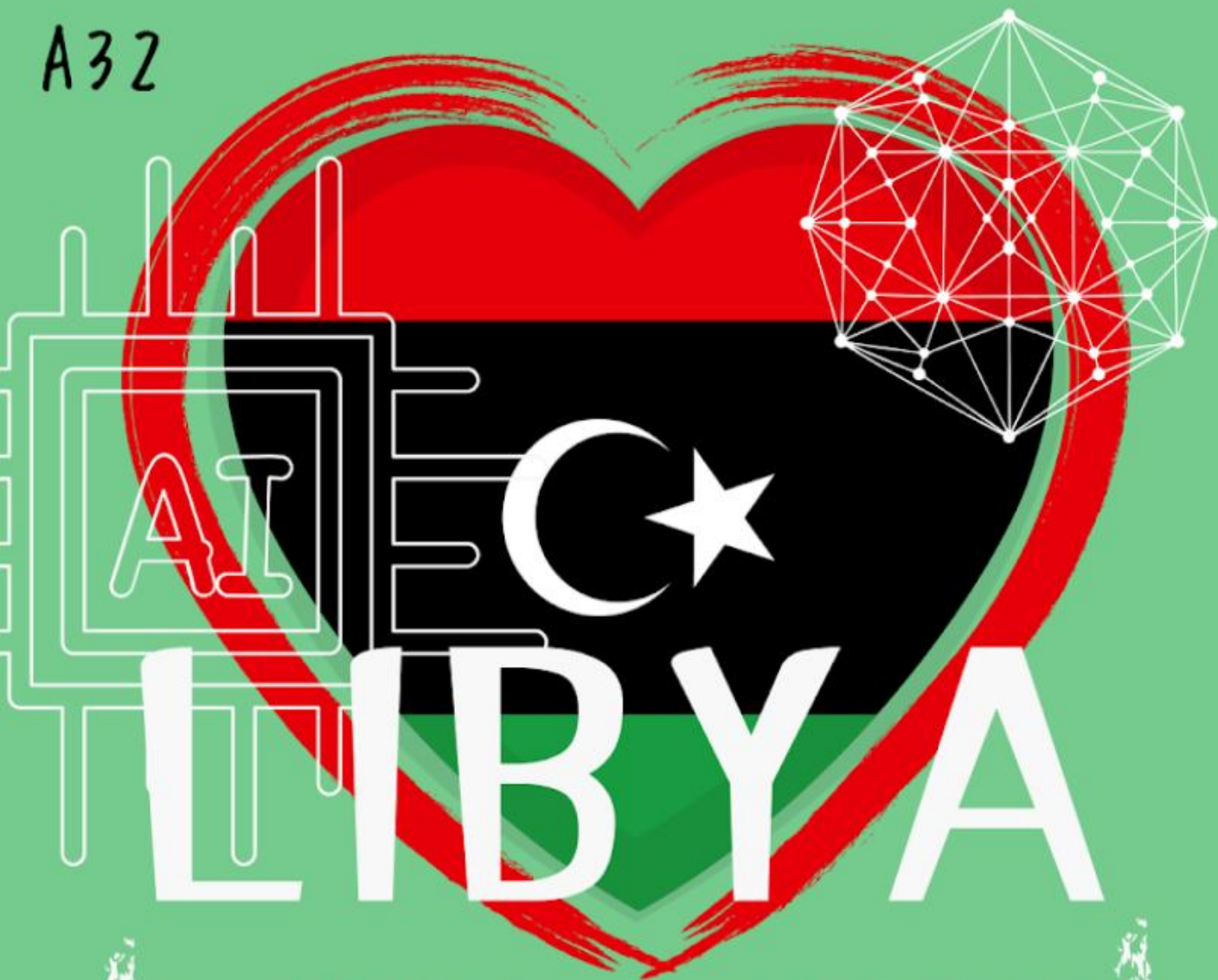
…人間中心のアプローチと国際人道法の遵守

論点②

…国際的な規範と倫理的ガイドラインの確立

我が国は、紛争の深刻な影響を直接経験してきた国として、「AI と軍事」という課題に対する国際社会の迅速かつ断固たる行動を強く求めます。

有意義な二日間にしましょう！よろしくお願い致します！



■論点1

AI兵器の全面禁止に反対し人間が命の判断を！

■論点2

AIによる誤作動の責任所在を明確にする国際的な基準の作成を！

私たちはAI兵器に慎重かつ現実的な姿勢を取ります。
会議では倫理と安全の両立を目指します。

対話を重ね互いの理解を深め、実りある会議にしていきたいと思います。

Lithuania



バルト三国の最南に位置するリトアニアです
2日間の会議、どうぞよろしくお願いします！

論点 1

- ・ 生成AIと軍事用AIの導入の推進
- ・ 軍事用AIが、国家間での不均衡を生まないための各国一律の規制、使用の促進

論点 2

- ・ 軍事用AIに関して「意味ある人間統制(MHC)」を国際法に盛り込む
- ・ 非致死用途を含む平和目的の生成AIの活用における国際協力の促進
- ・ 生成AIが提供する情報の透明性・公平性担保のために、出力内容の出典明記義務など国際的規範を設ける

Mali



論点1

他国への軍事用AIの規制を提案
軍事用AIの使用に関する世界共通の許可証
の作成、普及

論点2

・AIの開発者が責任をとるように提案
開発者には誤作動防止のための対策の開示、
またその際の責任の所在を承認する趣旨の文
書の提示を
義務付ける

2日間よろしくお願いいたします！



MEXICO

議題1 AIを搭載した軍事兵器に関する法的拘束力を持った国際文書を制定する会議を設ける。その際、人間の操作を逸脱するAI兵器を禁止し、また厳格な規制を課すことを提案する。

議題2 AI兵器により発生した国際法違反の結果について、各国家が最終的な責任を持つ、透明性の高い調査と被害者救済についての説明責任を含んだ枠組みが組み込まれた条約を作成すること。



当日は軍事利用に対して否定的な立場を取る予定です。
同じ考えを持つ国々、特にラテンアメリカ・カリブ海諸国の皆様とはぜひ積極的に議論を交わしたいと思います。
どうぞよろしくお願いします。

New Zealand

論点1

人間中心・倫理的アプローチを重視し、LAWS（致死的自律兵器）の**包括的規制**を求めます。

論点2

AI誤作動の責任は常に**人間**にあるべき。国際的な透明性と補償メカニズムが**不可欠**です。

有意義な会議に
しましょう！



NORTH KOREA

論点1

生成AIは全面的に規制を。
その中でもフェイクニュースや他国の情報
など真偽のつかないものへの規制強化
軍事用AIは安全保障・防衛面から必要！

論点2

責任はAIの開発者に。

使用した利用者に責任は帰属しない。

会議当日を楽しみにしています！

Legalize Allow Weapons to Train us

生成AIの使用：	個人情報保護、平等性・透明性担保、セキュリティ保護の上
軍事用AIの使用：	人への攻撃の最終判断がAIに委ねられている兵器の禁止
誤作動の対応：	企業が誤作動の可能性を通告していた→使用者責任 人為的誤作動→セキュリティ管理者の責任
誤作動の制裁議論：	ICJでの議論



A40

Oman



論点1

生成AIと軍事用AIのルール作り

論点2

AIが誤作動した時の責任と対応

中立的な立場に立ち、軍事におけるAIの利用に関する共通のルール作りを目指します！！中東の中でも平和なオマーンが世界と中東の橋渡し役となり、積極的に議論に参加します！

二日間有意義な会議にしましょう！





Pakistan

論点1

AI影響ケアプログラム (SAFE)

→AIの誤作動による被害のサポート

段階的適応規制モデル (GRADE)

→発展途上国の規制を調整する体制

論点2

共同責任制度の国際基準の確立

→リスクレベルを分類し、調整する

AI倫理・安全保障機関

→AI被害救済基金の運用



当日も宜しくお願い致します！
実りのある
会議にしましょう！



A43 PHILIPPINES

フィリピン大使です！

AIについて、「慎重な運用」「国際規範の遵守」
「人道的観点の重視」を基本方針としています。
ぜひ有意義な二日間にしましょう！

論点1

段階的な規制の導入、支援枠組みの併設
「意味のある人間の関与」

論点2

多重的责任モデル、責任主体の事前特定
事故・誤作動の調査を目的とした国際調査機関の創設



初めまして我々は**ポーランド**である
本会議では是非皆さんと協力することを望む

論点1

AIを管理している企業を政府が一時的に管理しAIに対する規制がなされているかを確認する。軍事用AI LAWSを一部規制することを目指す。

論点2

AIが誤作動したときへの対応としてはその誤作動を起こしたAIをすぐに全機能を停止させることができるプログラムを作成しておくこと。AIが誤作動がおこしたときの制裁を行ったりする新機関を設立することを目指す。

Portugal/



論点1 AIの使用に関する年齢制限を設けるべき

論点2 AIの誤作動を含め運用による利益・不利益を全て管轄する国際的な機関を設ける必要がある



Qatar



【論点1】

- ・ 軍用AIや非人道的生成AIの開発を規制
- ・ 利用企業、国家に情報公開を義務付ける
- ・ 高度知能や犯罪目的の技術にも規制を設置

【論点2】



- ・ AI誤作動の責任ルールを国際統一
- ・ 国際対応の法的ガイドライン策定
- ・ 軍事AIと現行ルールを各国と精査
- ・ 危険技術の開発を制限



AI推進国として、私たちは世界平和を目指し、
AIを世界に平和的かつ 有益に活用
するための取り組みを進めていき
ます。

**共にAIの良い使い方を考えて
いきましょう！**

A47

Russian Federation

各国の文化的価値観を第一に
共に最良の改善策を

生成AI

世界全体に、平等に
技術の共有を

LAWS

正確に定義し
適切な規制手段の選択を

2日間何卒よろしくお願いします

Serbia

論点1

責任の所在を明確にしたブラックボックスの義務化と法的拘束力のある措置についての議論



論点2

人間の関与を主とする国際規範の形成



セルビアは軍事的緊張を抱える地域に位置しているため軍事AIの利用について慎重に議論を進めたいです。2日間有意義な会議にできるよう頑張りましょう！！

Slovakia



論点 1

【生成AI】

① AIリテラシー教育の強化

ユネスコ主導で教育課程にAIリテラシーを組み込み、正しい理解と活用を促進

② ディープフェイクの生成・拡散の禁止

国家の分断や社会不安を招くディープフェイクの生成と拡散の国際的規制に向けた枠組み構築

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

【軍事用AI】

① LAWSの全面禁止

人間が一切関与しない兵器の開発・配備・使用を国際法で明確に禁止

② あくまで主体は人間

軍事用AIにおいて、「意味ある人間の関与」が無意味なものにならないよう、人間とAIの介入範囲の明確な線引きを行う

論点 2

【責任の所在】

- ・ **生成AI** → 誤作動に関わった主体が連帯して責任を負うべき
- ・ **軍事用AI** → 内部的・外部的誤作動の区別を明確にし、それに応じた責任の所在把握が必要

【国際的な基準】

- ① AIが結果を出力するまでのプロセス開示を義務化
- ② 責任主体に応じて、企業や個人は“国内法”、
国家は“国際司法裁判所”で裁かれるべき
- ③ 継続的に議論可能な 新たな国連専門機関 の創設を支持

曖昧な議論に終わらせることなく、具体的な合意形成を目指します
実質的な前進と成果のある議論を、一丸となって築いていきましょう！



SOMALIA

論点1

～「ヒューマン・イン・ザ・ループ」の義務化～

人間による最終判断を義務とする国際的規律を確立すること
途上国への支援体制を国連主導で議論すること

論点2

～AI誤作動時の対応の明確化～

被害国への補償体制の整備、責任の所在の構成
機関による事故調査の保証

各国大使へ

ソマリアは、長年にわたる内戦と政情不安により国家機能が弱体化し、
現在も多くの人々が貧困や暴力に直面しています。
一方で、豊かな海洋資源や若い人口を持ち、発展の可能性を秘めた国でもあります。
AIのルールづくりに参加し、国の課題も、世界の課題も、共に解決していきたい。
みなさん、2日間よろしくお願いします。

A51



SOUTH AFRICA

私たちは行動しなければいけません。

**AIによる線引きと分断
が生じるその前に。**

論点Ⅰ

差別を助長しない多角的データバンク

論点Ⅱ

**核のような持つ者・持たざる者の立場を生まない
LAWS完全規制**

2日間良い会議にしましょう！

South Korea



以下の内容を盛り込んだ
ガイドラインの作成を提案します！

論点1

- ・ 攻撃には人間の介入を必要とする
- ・ AIが作った画像などにAI生成マークをつける
- ・ 一部の分野に対して、AIに
出典が明確な情報のみを学習させる



論点2

- ・ AIが誤作動を起こした際の責任は国家または非国家
組織におけるAIの使用を指示した者に帰属する
- ・ 制裁に関する議論はGGEで行う
- ・ 各国政府は、民間人に対し適切なAIの導入方法や
活用方法について教育を受けられる場を提供する

自国は「意味ある人間の関与」を特に重視しています

1カ国でも多くの国の意見をDRに
盛り込むことをお約束します！



SOUTH SUDAN

論点1

AIの規制なし

論点2

国内に対しては責任を問わず、国外に対しては責任を利用している機関の最高責任者に委ねるものとする

よろしくをお願いします

Sudan A54

論点1

- Lawsの使用全面禁止
- 攻撃判断に人間の関与が必須

論点2

- 最終責任は国家元首、国防大臣が負う
- 民間AIの責任は開発者、企業に明確化



皆さんとの議論を
心待ちにしています！

Switzerland

01 論点 1

- ① 多言語ガイドラインを作る
- ② 軍事用AIの研究・分析を行い、必要なデータの公開を求める
- ③ UNODAにLAWSをどの国に輸出するべきかの基準を設ける組織を作る

02 論点 2

- ① AI技術の発展を妨げないように責任の所在は利用者と提供者にあるものとする
- ② AIによる国際的な被害についての制裁に関する議論をUNIDIRが行う

A56

Syria



論点 1

軍事用AIの使用をすべて禁止するのではなく、
軍事用AIを利用した無差別攻撃を禁止する

論点 2

AIの誤作動に対して制裁を課す
新たな規制機関の設立

長く内戦が続いているため、
軍事用AIの導入を容認しつつ、
人的被害の最小化を目指します。



* Tajikistan *

論点①

AIを使用するに至った経緯を明確にし、できる限りAIがどうしてその結論に至ったのかという過程を明らかにすることを約束する上で使用する
軍事用 AI を所有している国は所有していることを(国連に)表明する
また、戦争時にどのようなスタンスで軍事用 AI を使用するかも明確にする

論点②

軍事 AI の誤作動は国家が責任を負い、生成 AI に関しては開発者と提供者、利用者の共同責任とする
誤作動が起きてしまった場合には一時的な第三者委員会を設け、責任の所在を議論する
(常任理事国と非常任理事国が参加国)

皆さんと有意義な時間を過ごしたいです！！
よろしくお願いします！！！！



Thailand



論点1: AIに対する規制

- ・Human-in-the-loop(人間の関与)の国際原則化
→LAWSの使用・開発は禁止されるべき

- ・AIの信頼性を高める制度づくりが不可欠
→国家主導のAI透明性認証制度の創設
→軍事転用を防ぐ国際輸出管理枠組みの導入

論点2: AIの誤作動への対応

- ・AI誤作動の国際報告義務制度導入
→誤作動を速やかに共有・分析し再発を防ぐ
→国連が一元管理する中立的な報告システムを提案

- ・国際的な対応ガイドラインと責任保証枠組みの整備
→AI暴走時の停止手順・人道的対応の標準化
→被害者救済と国家・開発者の責任明確化

我々は中立性と地域協調を重視する立場から、
建設的な対話と合意形成に積極的に貢献する所存です。
どうぞよろしくお願いいたします。

A59

Turkey



論点 1

- ・ 地域裁量という形のガバナンス作成
- ・ **LAWS** の段階的禁止
- ・ 違反使用に対する国際的な枠組みの強化

論点 2

- ・ 第三者監査機関や影響評価、登録記録を義務づけ
- ・ **Ai** の誤作動による被害に対して報告を行う
- ・ **CCW** 会合や **GPAI**、 **OECD** の機関で
制裁や対応議論をすべき

2 日間、よろしくおお願いします。

頑張って国際合意を目指しましょう！

トルクメニスタン

論1)

生成AI：慎重的で段階的な導入

軍事AI：LAWsは全面的に反対

論2)

①ケースによって責任主体を吟味する

②①の線引きを公平に行う第三者機関の設立！

状況が異なる国々が平等な発言権を持つこの機会是真的持続的な平和に近づける最高の場です。

いい会議にしましょう！



United Arab Emirates

論点1

- ・生成AIの誤った使い方を防ぐガイドラインの作成

論点2

- ・誤作動が発生した状況の報告書の提出の義務化
- ・責任の所在は一般的には共同責任
 - 利用者の意図的な誤作動など、は利用者に責任を帰属

グループینگ

LAWSの「認識」「意志決定」「攻撃」三段階の運用をどの程度まで規制または容認するか

実り多い大会にしましょう
2日間よろしくお願いします！

USA

AIに力を。

人に「責任」を。

AIを生み出した私たち人類が、その活用について議論することは、私たちに課された「責任」です。

その「責任」をしっかりと果たすためにも、会議当日もどうぞよろしくお願いいたします。





Venezuela

論点1

完全自律型致死兵器の禁止と人間の
最終判断の保持を支持します。

論点2

人間による管理と透明な
検証制度の整備を提案します。



VIETNAM



論点①：

軍事用AIの透明性と人間による統制を義務づける国際的枠組みを導入する

論点②：

AI誤作動時の責任を明確にし、情報開示を義務づける国際制度を構築する。

軍事用AIに対して慎重かつ責任ある国際対応を求める立場を取ります。その実現に向けて、透明性・統制・ルール作りを各国とともに頑張っていきましょう。

**未来につながる二日間
にしていきましょう！**

