



UNITED NATIONS HIGH-LEVEL ADVISORY  
BODY ON ARTIFICIAL INTELLIGENCE  
(OFFICE OF THE SECRETARY-GENERAL'S  
ENVOY ON TECHNOLOGY)

---

---

— B A C K G R O U N D   G U I D E —

---

## Letter From The Executive Board

Distinguished Delegates,

It is with great pleasure that we welcome you to the United Nations High-Level Advisory Body on Artificial Intelligence at LMUN 2025. As we stand at the precipice of a technological revolution, the decisions we make today regarding artificial intelligence will fundamentally shape the future of humanity.

The two agendas before this committee address existential questions of our digital age:

How do we govern data that fuels AI while protecting individual privacy?

How do we harness AI's transformative potential while ensuring human dignity, accountability, and fundamental freedoms remain paramount?

Your role transcends that of ordinary delegates. The agendas before us are not theoretical, they are the core, immediate challenges facing the international community. Your role is not just to debate, but to innovate, forging a path toward AI that is safe, secure, and trustworthy, grounded in human rights and the pursuit of the public good.

We expect robust debate, innovative solutions, and most importantly, a commitment to multilateralism that transcends borders. In this model UN set-up the world is watching, and history will judge our ability to govern responsibly what may be humanity's most consequential invention.

With this, we welcome you once again to LMUN 2025 and wish you the best of luck with your preparation for this committee.

Warm regards,

Kinjalk Sharma  
Co-Chairperson, UNHLAB on AI  
LMUN 2025  
[unhlab.lmun2025@gmail.com](mailto:unhlab.lmun2025@gmail.com)

Aaryan Dhawan  
Co-Chairperson, UNHLAB on AI  
LMUN 2025  
[unhlab.lmun2025@gmail.com](mailto:unhlab.lmun2025@gmail.com)

Official email address :  
(for paperwork and queries)  
[unhlab.lmun2025@gmail.com](mailto:unhlab.lmun2025@gmail.com)

# ABOUT THE COMMITTEE

## Background and Mandate

The United Nations High-Level Advisory Body on Artificial Intelligence was established in October 2023 by UN Secretary-General António Guterres to examine the risks, opportunities, and international governance of artificial intelligence. This initiative emerged from growing recognition that AI technologies were developing at unprecedented speed, often outpacing existing regulatory frameworks and ethical guidelines.

The Advisory Body operates under the Office of the Secretary-General's Envoy on Technology and comprises 39 members representing governments, private sector, civil society, academia, and technical community from diverse geographical regions. Its mandate includes providing recommendations for international AI governance, identifying opportunities for AI to accelerate the Sustainable Development Goals (SDGs), and proposing frameworks for managing AI risks.

The Executive Board expects delegates to think creatively and practically. This is a policy-making body, not a legislative one. Resolutions should focus on governance frameworks, international norms, operational principles, and institutional mechanisms. We are looking for proposals on how the UN system and its member states can effectively manage and guide AI development.

### Specific Expectations:

- **Multi-stakeholder Perspective:** Consider the impact on governments, civil society, academia, and the private sector.
- **Actionable Policy:** Vague calls for "cooperation" are insufficient. Propose specific mechanisms (e.g., a proposed international data-sharing protocol, a mandate for an AI Risk Assessment Unit).
- **Nuance:** Recognize that AI is not a monolith. Solutions must be scalable, adaptable, and context-aware, balancing innovation with safety.

## **AGENDA 1: ADDRESSING THE GRAVITAS OF DATA GOVERNANCE FOR ARTIFICIAL INTELLIGENCE, INCLUDING INTERNATIONAL NORMS FOR DATA SHARING AND LEGALITY OF COLLECTING DATA FOR AI TRAINING ACROSS NEURAL NETWORKS WHILE PREVENTING OBSTRUCTION OF PRIVACY.**

### **Background and Context**

Data is the lifeblood of modern artificial intelligence systems. Machine learning algorithms, particularly deep neural networks, require vast quantities of data to identify patterns, make predictions, and generate outputs. This dependency creates a fundamental tension: the data necessary for AI advancement often contains personal information, copyrighted content, or sensitive material that raises significant legal, ethical, and privacy concerns.

The current landscape of AI data governance is fragmented and inadequate. Different jurisdictions have adopted varying approaches: from the European Union's comprehensive General Data Protection Regulation (GDPR) to more permissive frameworks in other regions. This fragmentation creates compliance challenges for global AI companies, regulatory arbitrage opportunities, and uneven protection for individuals' privacy rights.

Furthermore, the concentration of data in the hands of a few major technology companies creates power asymmetries. Nations and organizations lacking access to diverse, high- quality datasets face disadvantages in developing competitive AI systems. This "data divide" exacerbates existing inequalities between developed and developing nations, between well- resourced corporations and startups, and between different linguistic and cultural communities.

### **The Scope of the Challenge**

**Scale of Data Collection:** Modern AI systems are trained on datasets containing billions of parameters. Large language models may be trained on texts encompassing significant portions of the publicly accessible internet. Computer vision systems require millions of labeled images. Each data point potentially represents an individual's personal information, creative work, or private communication.

## **Types of Data Concerns:**

- **Personal Data:** Information that identifies or relates to individuals, including biometric data, browsing histories, social media activity, and location data
- **Copyrighted Content:** Books, articles, images, music, and code created by authors, artists, and developers
- **Sensitive Information:** Health records, financial data, political opinions, and other information requiring heightened protection
- **Technical Complexities:** The opacity of neural networks creates accountability challenges. Once trained on particular data, AI models embed patterns from that data in ways that are difficult to trace or reverse. This raises questions about the right to be forgotten, bias perpetuation, and the ability to audit AI systems for data misuse.

## **International Legal Frameworks**

### **Existing Instruments:**

- **GDPR (EU):** Establishes strict requirements for data collection, processing, and storage, including provisions for automated decision-making
- **CCPA (California):** Provides consumers rights regarding personal information collection and sale
- **Beijing AI Principles:** Emphasizes responsible data usage and privacy protection
- **OECD AI Principles:** Advocates for transparency and responsible stewardship of trustworthy AI
- **Gaps and Limitations:** Current frameworks were largely designed before the explosion of generative AI and large-scale training. They often fail to address:
  - Cross-border data flows for AI training purposes
  - The distinction between data collection for services versus AI training
  - Compensation mechanisms for creators whose work trains AI systems
  - Special protections for data from vulnerable populations
  - Standards for synthetic data and its relationship to real data

## **Critical Issues for Discussion**

### **1. Consent and Transparency**

Traditional consent models assume individuals can make informed decisions about data usage. However, when data is scraped from public internet sources or collected indirectly, meaningful consent becomes impractical. How can international norms balance innovation with genuine informed consent?

1. Questions to consider:

- I. Should AI companies be required to disclose all data sources used in training?
- II. What constitutes adequate notice when data will be used for AI training versus immediate service provision?
- III. How can consent mechanisms be designed for datasets containing billions of records?

### **2. Data Rights and Compensation**

Artists, writers, and creators increasingly find their work reproduced or imitated by AI systems trained on their content without permission or compensation. This raises fundamental questions about intellectual property in the age of AI.

• Questions to consider:

- I. Should creators have the right to opt out of AI training datasets?
- II. What compensation models might be appropriate when copyrighted work is used for commercial AI training?
- III. How do we balance fair use doctrines with creator rights in different legal traditions?

### **3. Cross-Border Data Governance**

AI development is inherently global, but data governance remains largely national. This creates jurisdictional conflicts and regulatory uncertainty.

Questions to consider:

1. What mechanisms could facilitate legal cross-border data sharing for AI research?
2. How can we prevent "data havens" with lax privacy protections from undermining global standards?
3. Should there be international standards for data localization requirements?

### **4. Privacy-Preserving Technologies**

Technical solutions like federated learning, differential privacy, and synthetic data generation offer potential pathways to train AI while protecting privacy. However, these techniques have limitations and trade-offs.

Questions to consider:

1. What role should privacy-enhancing technologies play in data governance frameworks?
2. How can we ensure these technologies are accessible to researchers and companies globally?
3. What standards are needed to verify that privacy-preserving claims are genuine?

### **5. The Data Divide**

Data availability varies dramatically across regions, languages, and communities. This affects whose perspectives are represented in AI systems and who can develop competitive AI technologies.

Questions to consider:

1. How can international cooperation facilitate equitable access to training data?
2. What mechanisms could support data collection and curation in underrepresented languages and regions?
3. Should there be data commons or shared repositories for AI research in the public interest?

## **6. Biometric and Sensitive Data**

Biometric data (facial recognition, voice patterns, gait analysis) and sensitive personal information present heightened risks when used in AI systems, potentially enabling mass surveillance or discrimination.

Questions to consider:

1. Should certain categories of data be prohibited from AI training entirely?
2. What special protections are needed for biometric data collection and use?
3. How do we balance security applications of biometric AI with privacy rights?

### **Key Questions to Consider**

1. How can international law adapt to address data collection practices specific to AI training, which differ fundamentally from traditional data processing?
2. What enforcement mechanisms could ensure compliance with international data governance norms across jurisdictions?
3. Should there be a distinction between data usage for AI systems serving the public interest (health research, climate modeling) versus commercial applications?
4. How can we protect privacy while enabling beneficial AI research that requires large datasets?
5. What role should international organizations play in facilitating data sharing agreements between nations?
6. How do we address the retroactive question of AI systems already trained on data collected without adequate consent or compensation?
7. What standards should govern the creation and use of synthetic data, and can it adequately substitute for real data?
8. How can indigenous communities and marginalized groups be empowered to control data about their cultures and identities?



## **AGENDA 2: ESTABLISHING ADEQUATE GUARDRAILS FOR THE REGULATION AND IMPLEMENTATION OF ARTIFICIAL INTELLIGENCE, WHILE PUTTING SPECIAL EMPHASIS ON HUMAN OVERSIGHT, ACCOUNTABILITY, LACK OF TRANSPARENCY, DIFFUSION OF RESPONSIBILITY, AND UPHOLDING FUNDAMENTAL FREEDOMS AT THE GLOBAL FORUM.**

### **Background and Context**

Artificial intelligence systems are being deployed across virtually every sector of society- healthcare diagnosis, criminal justice, financial services, education, employment, autonomous vehicles, military applications, and content moderation. While these applications promise significant benefits, they also pose unprecedented risks to human rights, safety, and democratic institutions.

The challenge of AI regulation is compounded by several factors: the rapid pace of technological change, the technical complexity that makes AI systems opaque even to experts, the global nature of AI development and deployment, and the tension between enabling innovation and preventing harm. Unlike previous technological revolutions, AI systems can make consequential decisions affecting millions of people in milliseconds, often with minimal human involvement.

Current regulatory approaches vary dramatically across jurisdictions. Some nations have adopted comprehensive frameworks, others rely on sectoral regulations, and many lack any AI-specific governance. This fragmentation creates uncertainty for developers, uneven protection for individuals, and the potential for regulatory arbitrage where companies seek jurisdictions with minimal oversight.

### **The Imperative for Human Oversight**

The Automation Paradox: As AI systems become more capable, there is a growing temptation to reduce human involvement in decision-making. However, this creates significant risks:

- **Deskilling:** Over-reliance on AI systems can erode human expertise and judgment
- **Automation Bias:** Humans tend to over-trust automated systems, accepting their outputs without adequate scrutiny
- **Accountability Gaps:** When decisions are delegated to AI, responsibility becomes diffused
- **Loss of Human Agency:** Fundamental life-affecting decisions may be made without genuine human consideration
- **Models of Human Oversight:**
  - **Human-in-the-Loop:** Humans make final decisions with AI assistance
  - **Human-on-the-Loop:** Humans supervise AI operations and can intervene
  - **Human-in-Command:** Humans set parameters and maintain ultimate control
  - **Meaningful Human Review:** Ensuring human involvement is substantive, not merely procedural

## **CRITICAL ISSUES FOR DISCUSSION**

### **1. Accountability Frameworks**

One of AI's most challenging aspects is the "responsibility gap", the difficulty of attributing accountability when multiple actors contribute to an AI system's development and deployment.

Questions to consider:

1. Who should be held accountable when an AI system causes harm: developers, deployers, users, or the organizations that own the system?
2. How can we ensure accountability when AI systems operate across multiple jurisdictions?
3. What liability frameworks are appropriate for different levels of AI autonomy?
4. Should there be strict liability for certain high-risk AI applications?
5. How do we handle cases where AI systems produce emergent behaviors not explicitly programmed?

### **2. Transparency and Explainability**

The "black box" nature of many AI systems, particularly deep neural networks, creates fundamental challenges for oversight, auditing, and contestability.

Questions to consider:

1. Should companies be required to disclose when AI systems are making decisions about individuals?
2. What level of technical transparency is necessary for different AI applications?
3. How can we balance commercial trade secrets with the public's right to understand AI systems affecting them?
4. What standards should govern "explainable AI" requirements for high-stakes decisions?

### **3. Diffusion of Responsibility**

Complex AI supply chains involving data providers, algorithm developers, cloud infrastructure providers, and deployment organizations create scenarios where responsibility is fragmented.

Questions to consider:

1. How can regulatory frameworks prevent organizations from obscuring responsibility through complex vendor relationships?
2. What documentation and traceability requirements should apply throughout the AI lifecycle?
3. Should there be mandatory impact assessments before deploying high-risk AI systems?
4. How can we ensure that responsibility cannot be evaded by claiming reliance on third-party components?

### **4. Fundamental Rights and Freedoms**

AI systems intersect with virtually every human right: privacy, freedom of expression, non-discrimination, due process, and autonomy. Ensuring AI respects these rights requires robust protections.

Non-Discrimination and Bias:

AI systems can perpetuate and amplify existing societal biases. Historical data often reflects discrimination based on race, gender, age, disability, and other protected characteristics. When AI systems are trained on such data, they can institutionalize bias at scale.

Questions to consider:

1. What standards should govern fairness in AI systems across different cultural contexts?
2. How can we audit AI systems for discriminatory outcomes?
3. Should certain applications (like predictive policing or automated hiring) face heightened scrutiny?
4. What remedies should be available to individuals harmed by biased AI decisions?

## 5. High-Risk Applications

Certain AI uses present particularly severe risks and may warrant prohibition or strict regulation:

**Autonomous Weapons Systems:** AI systems that can select and engage targets without human intervention raise profound ethical and legal questions.

Questions to consider:

1. Should fully autonomous weapons be prohibited under international humanitarian law?
2. What constitutes "meaningful human control" over weapons systems?

III. How can we verify compliance with restrictions on autonomous weapons?

**Social Scoring:** AI systems that comprehensively evaluate individuals' behavior and restrict their opportunities based on those evaluations threaten fundamental freedoms.

Questions to consider:

- I. Should social scoring systems be prohibited internationally?
- II. What distinctions exist between credit scoring and broader social scoring?
- III. How do we prevent private companies from implementing de facto social scoring?

**Predictive Policing and Criminal Justice:** Using AI to predict crime or assess recidivism risk can perpetuate discrimination and undermine presumption of innocence.

Questions to consider:

1. What safeguards are necessary when AI is used in criminal justice?
2. Should defendants have the right to challenge AI-based evidence and risk assessments?
3. How do we address bias in training data derived from discriminatory policing practices?

## 6. Cross-Sectoral Considerations

Healthcare: AI diagnosis and treatment recommendations affect life and death decisions. Errors can have catastrophic consequences, yet AI also promises to democratize access to quality healthcare.

Questions to consider:

1. What approval processes should govern medical AI systems?
2. How do we ensure human healthcare providers remain adequately involved?
3. What liability frameworks apply when AI systems contribute to medical errors?

Education: AI tutoring and assessment systems affect students' life trajectories. Automated systems may fail to recognize diverse learning styles or perpetuate educational inequities.

Questions to consider:

1. What oversight is necessary for AI systems evaluating student performance?
2. How do we ensure AI in education supports rather than replaces human educators?
3. What protections are necessary for student data used to train educational AI?

Employment: AI systems increasingly screen job applications, monitor worker productivity, and make termination recommendations.

Questions to consider:

1. Should workers have the right to human review of AI-based employment decisions?
2. What transparency requirements should apply to AI hiring systems?
3. How do we prevent AI-enabled worker surveillance from violating dignity?

Financial Services: AI credit scoring and fraud detection affect access to essential services.

Questions to consider:

1. What explainability requirements are necessary for consequential financial decisions?
2. How do we ensure AI doesn't discriminate in lending while allowing risk-based pricing?
3. What appeals processes should exist for AI-based financial determinations?

## **7. Enforcement and Compliance**

Regulations without effective enforcement are mere aspirations. AI governance requires robust monitoring, auditing, and penalty mechanisms.

Questions to consider:

1. What agencies or bodies should have authority to oversee AI systems?
2. How can international cooperation facilitate enforcement across borders?
3. What penalties are appropriate for violations of AI regulations?
4. Should there be certification requirements for high-risk AI systems?
5. How can we ensure resources for enforcement keep pace with AI deployment?

## **8. Innovation and Proportionality**

Excessive regulation could stifle beneficial AI development, particularly in under-resourced environments. Frameworks must balance protection with innovation.

Questions to consider:

1. How can regulations be risk-proportionate, applying lighter requirements to lower-risk systems?
2. What support should be provided to help smaller organizations comply with AI governance requirements?
3. Should there be regulatory sandboxes allowing controlled experimentation?
4. How do we ensure that regulations don't entrench advantages of established players?

## KEY QUESTIONS TO CONSIDER

1. What institutional mechanisms are needed at the international level to coordinate AI governance across nations?
2. How can we ensure that AI governance frameworks are adaptable to rapid technological change?
3. What balance should be struck between prescriptive rules and principles-based guidance?
4. How can developing nations be supported in building capacity for AI governance?
5. What role should civil society, academia, and the private sector play in AI governance beyond traditional government regulation?
6. Should there be an international registry of high-risk AI systems similar to clinical trial registries?
7. How can we ensure AI governance respects diverse cultural values while maintaining universal human rights standards?
8. What mechanisms could enable contestability—allowing individuals to challenge AI decisions affecting them?
9. How should we govern open-source AI systems where development is decentralized?
10. What sunset or review provisions should apply to AI regulations to ensure they remain relevant?

## **EXPECTATIONS FROM THE EXECUTIVE BOARD**

### **Before the Conference**

**Research and Preparation:** You are expected to thoroughly understand both agendas, your country's or stakeholder's position on AI governance, and the basics of the technical workings of AI systems to understand what you are dealing with. Familiarize yourself with key documents including the EU AI Act and UNESCO's AI ethics recommendation.

**Position Papers:** Submit a comprehensive position paper addressing both agendas from your assigned perspective. Your paper should demonstrate understanding of the issues, articulate your position clearly, and propose concrete policy solutions. The paper should not exceed the length of 1 A4-sized page, with the font style and size being Times New Roman and 12pt, respectively. The paper should be submitted on or before November 3, 2025 at [unhlab.lmun2025@gmail.com](mailto:unhlab.lmun2025@gmail.com). Please refrain from using generative text chatbots for the preparation of the position paper, as it would be looked upon negatively in the context of committee performance.

**Coalition Building:** Begin identifying potential allies and areas of common ground. AI governance requires multilateral cooperation use pre-conference communication to establish working relationships.



## **During the Conference**

**Substantive Engagement:** Move beyond rhetorical statements to engage with the technical and ethical complexities of the issues. Ask probing questions, challenge assumptions constructively, and build on others' ideas.

**Solution-Oriented Approach:** While articulating concerns and interests is important, prioritize developing actionable recommendations. The world needs concrete proposals, not just aspirational principles.

**Diplomatic Professionalism:** Maintain decorum, respect diverse perspectives, and seek common ground. The most effective diplomacy finds paths forward even amid disagreement.

**Evidence-Based Argumentation:** Support your positions with concrete examples, data, and logical reasoning. Reference specific AI incidents, research findings, and existing governance efforts.

**Creative Problem-Solving:** Many AI governance challenges lack obvious solutions. Be willing to propose innovative approaches, such as:

1. Novel institutional mechanisms
2. Hybrid public-private governance structures
3. Technical standards and certification regimes
4. Capacity-building initiatives
5. Phased implementation approaches

**Inclusive Representation:** Ensure that your solutions consider impacts on marginalized communities, developing nations, and future generations. AI governance must work for all of humanity.

# RECOMMENDED RESOURCES

## Essential Reading

1. EU AI Act: Official text and explanatory materials
2. UNESCO Recommendation on the Ethics of Artificial Intelligence (2021)
3. OECD AI Principles (2019, updated 2024)
4. UN General Assembly Resolution on AI (September 2024)

## Technical Understanding

5. AI Basics: Fundamental concepts of machine learning, neural networks, and generative AI
6. Privacy-Preserving Technologies: Federated learning, differential privacy, homomorphic encryption
7. AI Safety Research: Alignment problem, interpretability, robustness

## Rights and Ethics

8. Universal Declaration of Human Rights: Foundation for rights-based AI governance
9. International Covenant on Civil and Political Rights: Relevant to AI surveillance and profiling
10. Academic Literature: Papers on AI ethics, fairness, accountability, and transparency

## Regional Perspectives

11. China's AI Regulations: Understanding diverse governance approaches
12. African Union AI Strategy: Representing Global South perspectives
13. ASEAN Guide on AI Governance: Regional cooperation models

## CLOSING THOUGHTS

The governance of artificial intelligence represents one of the defining challenges of our era. The decisions made in forums like this will shape whether AI becomes a tool for human flourishing or a source of unprecedented harm.

Your participation in this Model UN is more than an academic exercise. It is preparation for the real work of global cooperation that your generation will inherit. The dilemmas we explore here- balancing innovation with protection, respecting privacy while enabling progress, ensuring accountability in complex systems, are not hypothetical. They are urgent, consequential, and require your most thoughtful engagement.

Approach these agendas with the seriousness they deserve. Question assumptions, seek evidence, consider multiple perspectives, and craft solutions that can actually work. The world needs not just critics of AI, but architects of governance frameworks that can guide humanity through this transformation.

The future is not predetermined. It will be shaped by the choices we make today and the frameworks we establish now. Make them count. See you at LMUN. Vive La Martiniere!

---



# THE TWELFTH SESSION