



**WSIS
FORUM 2026**

**6-10 July 2026
Geneva, Switzerland**

Session Outcome Document

Anchoring AI Evolution in Humanity: Pathways for Ethical, Inclusive, and Cross-Boundary Collaboration

Zhu Consulting

Monday, 6 July 2026, 11:00 - 11:45 AM CEST

<https://www.itu.int/net4/wsis/forum/2026/Agenda/Session/191>

Key Issues discussed:

- Three panelists attended, including:
 - Ms. Hao Ji Zhu, Founder, [Zhu Consulting](#) (moderator)
 - Ms. Yu Ping Chan, Head of Partnerships and Engagement, [UNDP Digital, AI and Innovation Hub](#) & Coordinator, [UNDP-Singapore Global Center](#)
 - Dr. Peter Slattery, Research Scientist, MIT FutureTech, [MIT AI Risk Initiative](#)
- Discussion examined current challenges in responsible AI implementation, high-leverage pathways for near-term principle implementation, and strategies for more effective cross-boundary ecosystem partnerships.

Key Outcomes of the session:

- **Imminent catastrophic risks and the responsibility gap:** New MIT AI Risk Initiative data reveals 18 of 24 risk domains carry $\geq 10\%$ probability of catastrophic outcomes within five years. Top harms include dangerous capabilities, competitive dynamics, weapons and cyberattacks, power centralization, and false information. While public users are most vulnerable, the primary burden of mitigation falls on developers and governance actors operating under an "illusion of responsibility," assuming others act. This is worsened by a lack of systemic understanding delineating power and accountability, exacerbated by the absence of shared taxonomies, consolidated evidence, and interoperable frameworks. Effective risk management demands substantial mitigation and structural coordination, including shared knowledge infrastructures, yet most institutions and countries lag behind.
- **The asymmetry of the Global South:** Risks are unequally distributed. Developing countries face amplified challenges due to infrastructure (e.g. electricity and internet connectivity) and governance capacity gaps, deepening the digital divide. Imported models and regulatory standards often fail to address local realities. With power concentrated in the Global North, developing nations in the Global South lack the leverage to "vote with wallets" like their Northern counterparts. The disproportionate impact of AI on the Global South requires further studies.

- **Narrow definition of capacity building:** Traditional rapid training is insufficient for building the true institutional readiness and resilience required for responsible AI adoption and governance. There is a critical need to shift from isolated knowledge transfer to problem-specific and long-term institutional adaption.
- **Coordination deficit:** The lack of incentives for coordination is a key failure point in ecosystem collaboration. Addressing this requires a shift in donor and institutional metrics to value and fund the "invisible work" of multistakeholder alignment.

Key Recommendations and Forward-Looking Actions:

- **Humanity-centric framing:** Anchor on the purpose of humanity, treating responsible AI as a continuous negotiation, with the protection of the most vulnerable as a guiding imperative.
- **"Problem-first" AI adoption:** Shift from blind adoption of existing technologies to a demand-driven, context-specific, and people-centered approach, ensuring AI solutions address the tangible realities and needs of local communities.
- **"Small AI" and localized alternatives:** Pivot to "small AI" and open-source models tailored for local contexts as viable, more cost-effective, and more locally appropriate alternatives for developing countries, reducing dependency and concentration.
- **Creating enabling environment through localized capacity building:** Cultivate enabling environment via localized capacity building, including expanding efforts beyond standalone training to long-term context-specific institutional adaptation. Mobilize diverse cross-sector partnerships to enable this type of assistance.
- **Systems mapping and targeted accountability:** Formulate granular systems mapping to delineate roles, incentives, and barriers across all levels, distinguishing between lack of will and lack of guidance. Adopt a co-design process where broadly applicable principles are contextualized locally, allowing local contexts, especially in the Global South, to feed insights into broader framework developments while preserving interoperability. Expanding empirical repositories allows us to replace vague mandates with documented actionable mitigations and strengthens local agency.
- **Aligning incentives via coordinated legal, civic, and market pressures:** Adopt a triad of legal, civic, and market pressures to realign incentives toward safety, especially through advocacy and regulation. Layered interventions can create stronger momentum to make irresponsible behavior costly and responsible adoption the optimal choice with the least resistance.
- **Performance metrics and funding models for coordination:** Revise performance metrics and funding models to actively reward cross-silo collaboration,

recognizing that partnerships is a prerequisite for translating high-level principles into the reality. By allocating resources to mechanisms and initiatives that facilitate dialogues and joint actions across sectors, regions, and levels, the ecosystem can shift into a more unified and efficient response to the complexities of AI risk.

- **Coalitions and UN's role as a coordinator:** To counter market asymmetry, Global South developing countries could benefit from collective capacity, aggregating isolated capabilities and efforts into strategic coalitions with leverage and ability to negotiate better terms and exert stronger influence on global norm-setting. UN can support this coordination, facilitating an inclusive global AI architecture.