

AI Standards for Global Impact: From Governance to Action

2025 Report

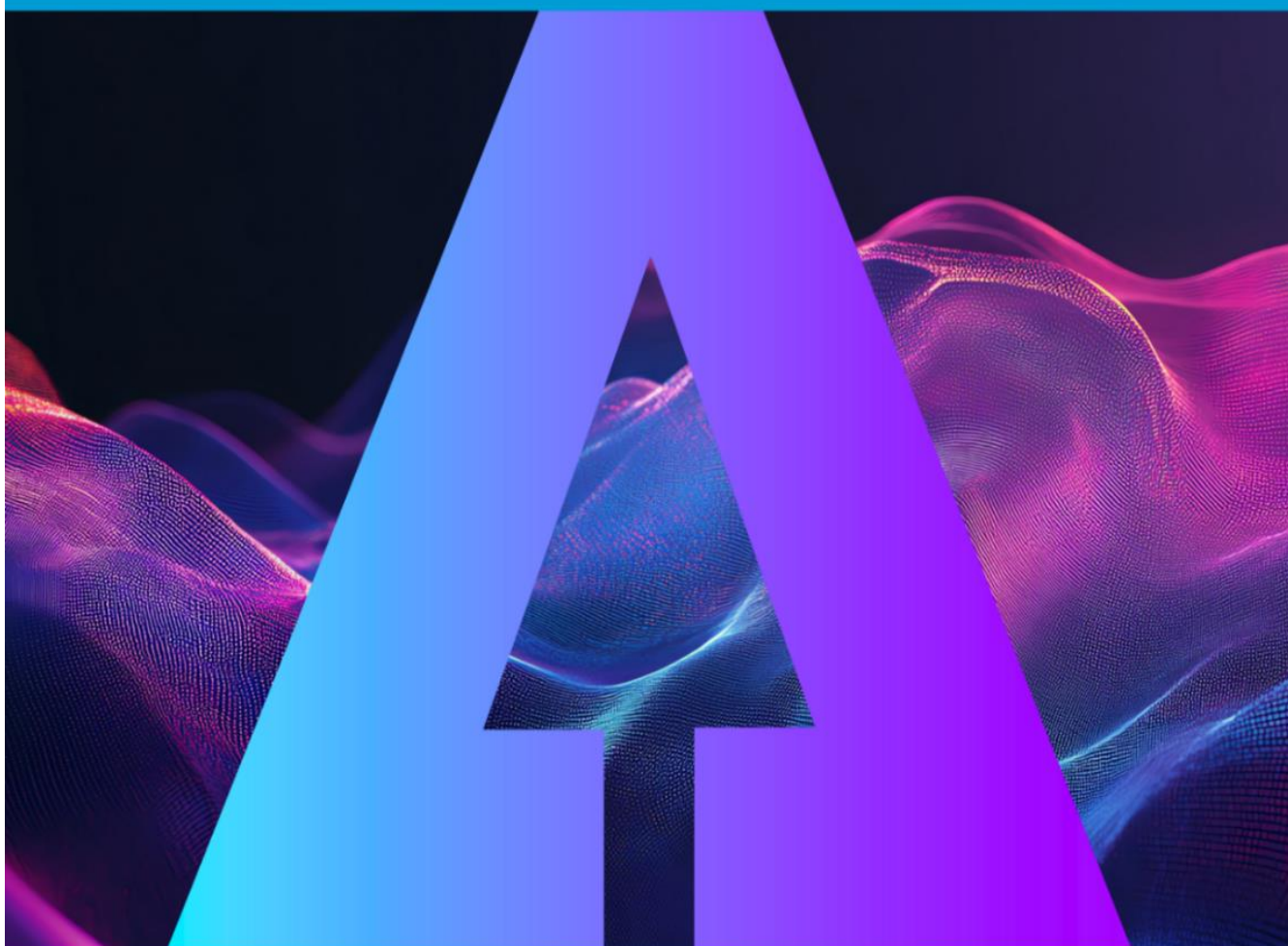


Table of Contents

Foreword	7
About AI for Good Global Summit 2025	8
Executive Summary	10
Part 1: International AI Standards Exchange	12
1 Importance of AI standards	13
1.1 AI standards at play	13
1.2 AI standards role and impact	13
2 AI standards exchange database	15
3 Role of International standards in AI	17
4 High-level roundtable - From principles to practice: how AI standards enable effective governance 22	
5 High level panels and keynotes	28
5.1 Transforming telecoms with AI and machine learning	28
5.2 AI Standards: Building trust and supporting innovation in a networked world	31
5.3 Navigating tomorrow: Education, skills, and standards in an AI-driven workplace	32
5.4 Computing and AI: Endless frontier and exploration	34
5.5 AI and multimedia authenticity	36
5.6 Humans and AI: The Journey	38
5.7 Donate your brainwaves for social impact	40
5.8 AI for food systems	42
5.9 Future of smart mobility	43
5.10 AI and energy	45
5.11 Elevating AI skills for all	47
Part 2: Thematic AI Standards Workshops	50
6 Trustworthy AI testing and validation	51
6.1 AI system testing	51
6.2 International collaboration on AI testing	53
7 Open dialogue on trustworthy AI testing	57
7.1 Outcomes	58
7.1.1 Group 1: Capacity building	58
7.1.2 Group 2: Standards, Best Practices and Conformity Assessment	58
7.1.3 Group 3: Institutional frameworks	59
7.2 Future directions	59
8 Enabling AI for health innovation and access	59
8.1 Global Initiative on AI for Health (GI-AI4H)	60
8.2 Strategy and Frameworks for AI in health	61

8.3	Use cases of AI in health	62
8.4	IP management to enable AI implementation for healthcare innovation and access	63
8.5	Technical brief on AI in traditional medicine	65
8.6	Closing summary	65
9	Human-centred AI for disaster management: empowering communities through standards	66
9.1	AI for disaster management	66
9.2	Global collaboration on standards for disaster management	67
9.3	Use cases of AI for disaster management	67
9.4	How AI transforms disaster management	68
9.5	Conclusion.....	69
10	AI and multimedia authenticity standards.....	70
10.1	Overview of AI and Multimedia Authenticity Standards Collaboration	70
10.2	Standards mapping landscape on AI and multimedia authenticity.....	72
10.3	Policy Paper: Building trust in multimedia authenticity through international standards	74
10.4	Fighting misinformation through fact-checking and deepfake detection.....	79
11	AI and machine learning in communications workshop.....	82
11.1	Introduction	82
11.2	Innovations in AI models	83
11.3	Standards and open source.....	84
11.4	Outcomes	85
11.5	Tools and Simulators, and Datasets.....	86
11.6	Opportunities for networking standards	87
12	Challenging the status quo of AI security	88
12.1	Keynote: Framing agentic AI and identity with a strategic lens	89
12.2	Agentic AI security	89
12.3	Key takeaways.....	95
12.4	Agentic AI identity management.....	96
12.5	Interplay of AI and cybersecurity: The good, the bad, and the ugly	97
12.6	Next steps	99
13	Empowering innovative and intelligent solutions at the edge (edge AI)	100
13.1	What is edge AI and why is it becoming important	100
13.2	Where is edge AI needed?	101
13.3	Edge AI hardware	102
13.4	Edge AI use case implementations	103
13.4.1	Addressing untreated hearing loss	103
13.4.2	MountAIn Edge AI infrastructure in remote environments	103
13.4.3	Building trustworthy local AI for healthcare at the edge	104
13.4.4	Implementing edge AI in mission-critical industries.....	104
13.4.5	From classroom to community: Five years of TinyML academic network impact	105

13.4.6	AI inference chip	105
13.4.7	Day 0 support for novel models enabling instantaneous deployment of edge AI	106
13.4.8	Defining the boundaries of AI from fundamental limitations to resource constraints ...	106
13.5	Key takeaways	106
14	Navigating the intersect of AI, environment and energy for a sustainable future	107
14.1	Understanding AI's environmental impact	107
14.2	Innovations in environmentally efficient AI	110
14.3	AI-driven environmental sustainability across industries	112
14.4	Standards, policies and regulations for sustainable AI and environment	114
15	Future Networked Car Symposium 2025	115
15.1	How AI can improve mobility for a better future	116
15.2	The revolution of AI automotive systems infrastructure	116
15.3	AI automotive applications ethics and their human elements	116
15.4	MoU between ITU and SAE	116
15.5	The role of standards in shaping autonomous mobility	117
16	AI readiness workshop	118
16.1	Introduction	118
16.2	Pilot AI readiness plugfest presentation and partner sharing	119
16.3	Outcomes	120
17	AI for agriculture – shaping standards for smart food systems	120
17.1	Dialogue on digital agriculture roadmap	121
17.2	Demo: Project Resilience	121
17.3	Farming the future – Laying the groundwork for scalable AI	122
17.4	Key outcomes	123
18	Women leaders in AI and standards	124
Part 3: Future AI standards for frontier technologies		127
19	AI and virtual worlds: Building the cities and governments of tomorrow	128
19.1	Virtual worlds, metaverse and citiverse – need for standards	128
19.2	Policy frameworks for virtual worlds	130
20	Brain-computer interfaces	130
20.1	Keynote address	131
20.2	Opportunities and challenges of BCI	131
20.3	Industrialization and commercialization barriers of BCI	132
20.4	Standardization and global cooperation	133
20.5	Key outcomes	133
21	Building the technical foundations for embodied intelligence in connected ICT environments	134
21.1	Defining embodied intelligence	134
21.2	Technical and infrastructural challenges	134
21.3	Standardization gaps and needs	135

21.4	Collaboration across standards bodies.....	135
21.5	Looking forward.....	137
Part 4: Quantum for Good.....		138
22	Quantum for Good Track.....	139
22.1	From promise to progress: Focusing on real use-cases.....	140
22.2	Quantum dilemma and risks	142
22.3	Inclusion and workforce: Building capacity for all.....	142
22.4	International collaboration and standards for quantum	143
22.5	The road to lasting impact	145
22.6	Future directions	146
Acknowledgements.....		147

Foreword

Standards have always been central to the work of ITU. Artificial intelligence reminds us why. The principles we want to live by must be embedded in the standards we develop.

Standards help us innovate at scale and share innovation worldwide. They are the common understandings we all need to prosper. Ultimately, standards create confidence to keep innovating, investing, and doing business internationally.

Standards must deliver this value far and wide, especially as we develop and apply advanced technologies like AI and quantum.

This report covers insights shared as part of the International AI Standards Exchange track of the AI for Good Global Summit 2025. It highlights the latest trends and applications in areas from networking, multimedia authenticity, cybersecurity, and quantum technologies to energy efficiency, healthcare, food security, and smart mobility.

In these areas and many more, the report summarizes the summit's in-depth workshop discussions on key innovations and related business and policy objectives, identifying related priorities for ongoing or new standards work.

To achieve the future we want, standards need to be the result of a process that is inclusive, transparent, and aligned with our ambitions for a better world.

ITU offers this assurance.

All participants' voices are heard, and every step forward is determined by consensus decision. We have 160 years of experience to build on. And we have global community that believes in commitment to collaboration and consensus.

I welcome you to join us.

Seizo One

Director of ITU Telecommunication Standardization Bureau



Figure 1: Seizo Onoe, Director of ITU Telecommunication Standardization Bureau

AI for Good Global Summit 2025

The AI for Good Global Summit took place from 8 to 11 July 2025 at the Palexpo International Conference and Exhibition Centre in Geneva, Switzerland. The mission of AI for Good is to unlock AI's potential to serve humanity through building skills, AI standards, and advancing partnerships. As the leading UN platform for dialogue on Artificial Intelligence (AI), the Summit is organized by the International Telecommunication Union (ITU) in partnership with 53 UN agencies and co-convened with the Government of Switzerland.



Figure 2: Doreen Bogdan-Martin, Secretary General, ITU



Figure 3: Tomas Lamanauskas, Deputy Secretary-General, ITU

The event was attended by over 11,000 delegates onsite from 170 countries, including Heads of UN agencies, over 100 government leaders, and distinguished delegates from industry, standards development bodies, academia, and civil society.



Figure 4: AI for Good Global Summit 2025 centre stage

Executive Summary

The AI for Good Global Summit 2025 featured a day dedicated to standardization on 11 July 2025 titled the "International AI Standards Exchange," organized by the International Telecommunication Union (ITU) in partnership with the International Organization for Standardization (ISO) and the International Electrotechnical Commission (IEC).

The International AI Standards Exchange and supporting workshops from 8 to 11 July examined the latest developments in AI and supporting standards development to galvanize global collaboration on AI for good.

The summit's programme featured the following standards-focused segments:

- a) High-level AI Standards Panel and launch of AI Standards Exchange Database on the Centre Stage on 11 July 2025
- b) High-level keynotes on Centre Stage on 11 July 2025
- c) High-level roundtable - From principles to practice: How AI standards enable effective governance
- d) A series of 13 thematic workshops from 8 to 11 July 2025 featuring standards work related to AI in different sectors and industry verticals such as:
 - i. [AI for agriculture – Shaping standards for smart food systems](#)
 - ii. [Enabling AI for health innovation and access](#)
 - iii. [Human-centred AI for disaster management – empowering communities through standards](#)
 - iv. [Navigating the intersect of AI, environment and energy for a sustainable future](#)
 - v. [AI readiness – Towards a standardized AI readiness framework](#)
 - vi. [Trustworthy AI testing and validation](#)
 - vii. [Open dialogue for trustworthy AI testing](#)
 - viii. [AI and machine learning in communication networks](#)
 - ix. [AI and Multimedia authenticity standards](#)
 - x. [Future networked car symposium 2025](#)
 - xi. [Empowering innovative and intelligent solutions at the edge](#) (Edge AI)
 - xii. [Challenging the status quo of AI security](#)
 - xiii. [Women leaders in AI and standards](#)
- e) To complement the AI focus and in celebration of the International Year of Quantum 2025, a dedicated Quantum for Good track now explores how quantum technologies intersect with global priorities, highlighting industry leadership and real-world innovation, the role of quantum in diplomacy, access and inclusion, the broader societal impact of quantum advances, and related standardization dynamics.
- f) Future standards for frontier technologies - Workshops from 8 to 10 July and sessions on the Centre Stage dedicated to emerging technologies and areas for future AI standards work:
 - i. [Brain computer interface \(BCI\) – Key technical standards and diverse application scenarios](#)

- ii. [Building the Technical Foundations for Embodied Intelligence in Connected ICT Environments](#)
- iii. [AI, virtual worlds & the human-centric citiverse- Building the cities and governments of tomorrow](#)

Key announcements made at the summit:

1. The new [Global Initiative on AI for Food Systems](#) led by ITU, the UN Food and Agriculture Organization (FAO), the World Food Programme (WFP) and the International Fund for Agricultural Development (IFAD) to use AI to boost productivity, efficiency, and global food security.
2. Two [landmark resources](#) on standards and policy considerations for multimedia authenticity were released by the [AI and Multimedia Authenticity Standardization Collaboration](#) driven by ITU, International Organization for Standardization (ISO), International Electrotechnical Commission (IEC) and other key standards communities. The collaboration is advancing standards to detect deepfakes and verify multimedia authenticity and provenance.
3. A new [AI Standards Exchange Database](#) to support cohesive standards development and application. The database currently includes standards from ITU, ISO, IEC, the Institute of Electrical and Electronics Engineers (IEEE), and the Internet Engineering Task Force (IETF), and welcomes contributions from all AI standards communities. The database will help standards developers to coordinate work and stakeholders to apply comprehensive standards suites.
4. ITU, ISO and IEC will bring together standards developers and stakeholders at the [International AI Standards Summit](#) from 2 to 3 December in the Republic of Korea.

Part 1 : International AI Standards Exchange

1 Importance of AI standards

1.1 AI standards at play

The development and adoption of AI standards will be essential for ensuring that AI technologies are beneficial for all. AI standards can provide a robust framework that guides the responsible development and deployment of AI systems with established best practices, fostering trust, innovation, and compliance. The continuous advancement of AI promises to transform nearly every sector of the economy.

As the AI landscape continues to evolve, standards will need to adapt to new challenges and innovations. Understanding them and their impact is essential for anyone involved in the development, deployment, or governance of AI technologies.

AI standards are **voluntary**, but they play a key role in the adoption of **good practices on trustworthy AI**, and effective technical solutions and methods to frame the deployment of AI systems. Among other things, they can facilitate compliance with legislation, promote innovation and ensure access to markets.

Some of the benefits of AI standards are

- Helping ensure the protection of fundamental rights, security, and transparency.
- Fostering innovation by establishing clear technical requirements.
- Enabling competition and market entry, especially for start-ups and SMEs.
- Supporting harmonized regulation and the interoperability of AI systems.

1.2 AI standards role and impact

AI standards can play an essential role in development and use of AI systems, for example:

- Defining methods for testing and evaluating AI systems
- Establishing performance criteria and technical requirements
- Facilitating risk management and regulatory compliance
- Supporting global trade
- Optimizing energy use
- Promoting societal welfare and enhancing public trust

AI standards can help ensure that AI systems meet predefined criteria for safety, reliability, and ethical behaviour. This instils confidence among users and stakeholders, fostering trust in AI technologies.

AI standards can help identify and mitigate risks associated with AI deployment, contributing to better risk management practices and clearer liability allocation. By providing guidelines for risk assessment and management, standards can help organizations identify potential risks

early and implement measures to mitigate them. This reduces the likelihood of adverse outcomes and clarifies the allocation of liability in case of failures or incidents involving AI systems.

International standards for AI are valuable tools both for regulators and companies, as they are developed based on market needs and help to facilitate global trade. Governments around the world are increasingly introducing policies and regulations to govern AI and related digital technologies, yet they are still at early stages of development. They also tend to be sector-specific, and their stringency varies from sector to sector. International standards are seen as a means to facilitate global trade because they can help harmonize requirements across countries, reduce duplication, and provide transparency.

The dual role of AI – as a significant consumer of energy and a contributor to GHG emissions – highlights the importance of integrating sustainability into AI development and deployment. Standards will be crucial in guiding this integration, helping ensure that AI systems are not only innovative but also environmentally responsible. By addressing key areas such as product efficiency, site operations, network performance, and resource management, standards can lay the groundwork for reducing the environmental footprint of AI. Concurrently, standards can also enable the strategic application of AI to support and accelerate climate goals, helping to optimize energy efficiency and enhance resilience through advanced solutions.

Standards development that includes ethical considerations and human rights can help ensure that AI systems are developed and deployed in ways that respect human dignity and societal values. This, in turn, enhances public trust and acceptance of AI technologies, paving the way for broader adoption and positive societal impact.

2 AI standards exchange database

The [AI Standards Exchange Database](https://aiforgood.itu.int/ai-standards-exchange/) supports the United Nations Global Digital Compact (Articles 58, 59 and 60), which proposed the establishment of an AI Standards Exchange to facilitate sharing of information about AI standards. At the first International AI Standards Summit held in India in October 2024 alongside WTSA-24TSA), the World Standards Cooperation (IEC, ISO, and ITU), proposed to set up the AI Standards Exchange Database by 2025 in response to the Global Digital Compact. The AI Standards Exchange Database was set up by ITU under the World Standards Cooperation and in partnership with the Institute of Electrical and Electronics Engineers (IEEE), the Internet Engineering Task Force (IETF), and the AI Standards Hub.

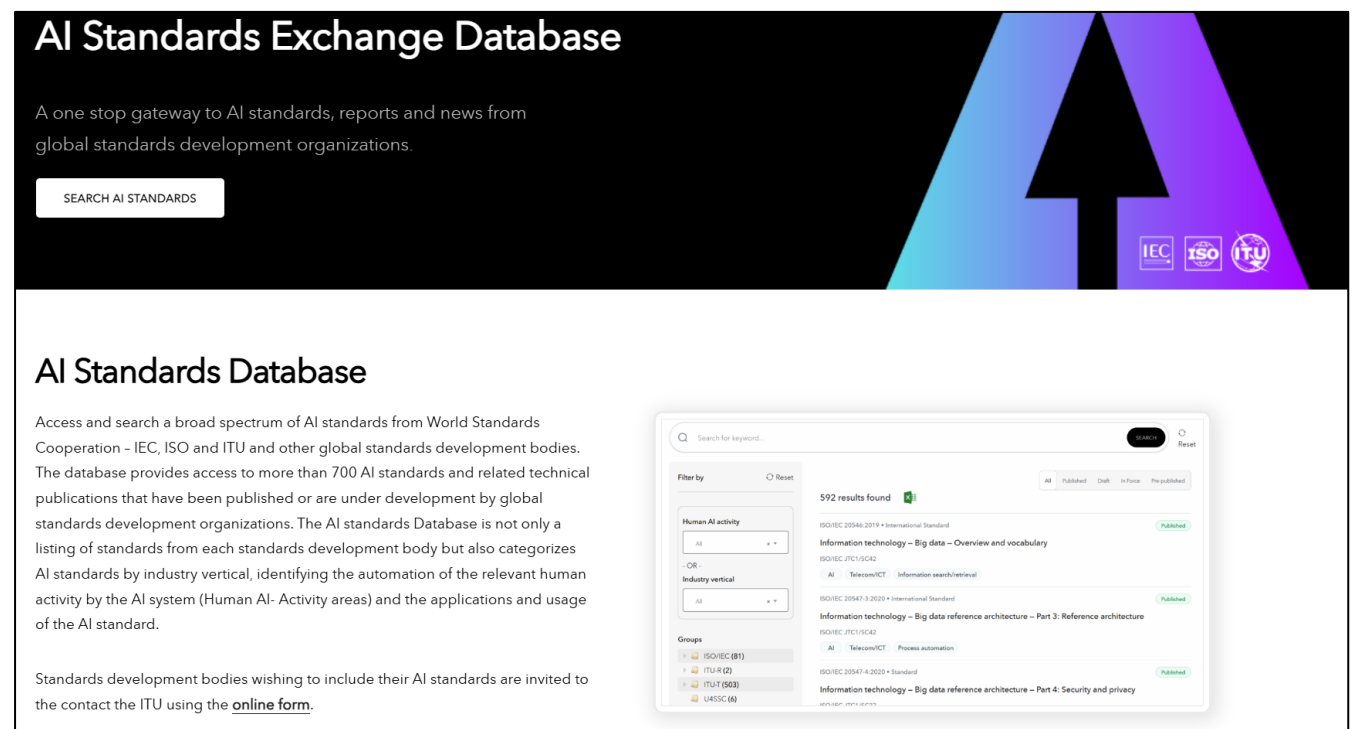


Figure 5: AI Standards Exchange Database¹

What is the AI Standards Exchange Database?

The [AI Standards Exchange Database](https://aiforgood.itu.int/ai-standards-exchange/) currently provides access to over 700 AI standards and technical reports from IEC, ISO, ITU, IEEE, and IETF.

The database provides a one-stop gateway with advanced search and filtering tools, offering access to AI standards classified by industry verticals, human-AI activities (i.e. the human activity that is automated by the AI system), applications, and use cases. It tracks both in force and draft (under development) AI standards. The database supports technical cohesion and interoperable solutions – key aims of the Global Digital Compact adopted last year as part of the United Nations Pact for the Future.

¹ See: AI Standards Exchange Database - <https://aiforgood.itu.int/ai-standards-exchange/>

The following metadata is used for AI-related publications and can be grouped under three main categories:

- a) General information about the AI publication (i.e. title, standards body issuing the publication and relevant technical committee, abstract, type of AI publication, date published, status, URL, etc.)
- b) AI use (industry vertical, human-AI activity, application area, use case description and URL)
- c) Other attributes (e.g. open source (Yes/No), standard/publication, open-source code repository, etc.)

The AI Standards Exchange Database classifies the type of AI publications as follows:

- 1. International standards
- 2. National standards
- 3. Technical reports
- 4. Technical specifications
- 5. Technical papers
- 6. Guides
- 7. Other publications

The AI use category metadata classifies AI publications further according to the following metadata: industry vertical, human AI activity, application area of interest, and information on use cases implemented (with URL).

The AI publications themselves are not stored in the AI Standards Exchange Database. The AI Standards Exchange Database provides a link to the webpage where a publication can be downloaded or more information about the publication's availability can be found (e.g. online purchase or other means).

A Coordination Committee composed of the representatives of standards bodies contributing to the databases will be established by ITU to oversee the database's governance and content updates, metadata fields required, and ensuring continuous contributions and momentum from global stakeholders. The AI Standards Exchange Database will be extended in future to include standards-based capacity-building initiatives.

The AI Standards Exchange Database welcomes input on new AI standards and technical reports from all standards bodies. Standards bodies interested in submitting their AI standards and technical reports for publication on the database can submit the information about their AI publication via the online form on the AI Standards Exchange database website.

3 Role of International standards in AI

The AI Standards Exchange Database was announced and showcased on 11 July 2025 at the AI for Good Global Summit during a High-Level Panel on the *Role of International Standards in AI*.

The High-Level Panel was moderated by Bilel Jamoussi, Deputy Director of the ITU Telecommunication Standardization Bureau, and the panelists were:

- Amandeep Singh Gill, Under-Secretary-General and Special Envoy for Digital and Emerging Technologies, Office for Digital and Emerging Technologies, United Nations
- Sung Hwan Cho, President, ISO
- Seizo Onoe, Director of the ITU Telecommunication Standardization Bureau
- Philippe Metzger, Secretary-General and CEO, IEC
- Kathleen A. Kramer, President and CEO, IEEE
- Paul Gaskell, Deputy Director of Digital Trade, Internet Governance & Digital Standards at the Department for Science, Innovation & Technology, UK



Figure 6: Left to right: Bilel Jamoussi, Amandeep Singh Gill, Sung Hwan Cho, Seizo Onoe, Kathleen A. Kramer, Philippe Metzger, Paul Gaskell

Panelists highlighted the importance of inclusive standards processes, building AI standards capacity around the world, strong connections among standards developers and their growing range of stakeholders, and promoting a human-centred AI ecosystem where standards address societal as well as technical challenges.

The AI standards exchange database forms part of broader collaboration to ensure standards provide practical tools to better shape AI. It will help standards bodies to coordinate their work

and empower companies, policymakers, and regulators with comprehensive suites of AI standards.

To ensure clarity, coherence, and impactful AI standards, standards developers should focus on four key elements:

- Translation (turning principles into practical standards through collaboration)
- Structure (a systematic approach to capacity building and quality digital infrastructure)
- Inclusion (broad participation and equal voices in decision-making)
- Connection (engaging diverse stakeholders, governments, and organizations to address real-world needs)

Amandeep Singh Gill, Under-Secretary-General and Special Envoy for Digital and Emerging Technologies, Office for Digital and Emerging Technologies, United Nations – *"We need to connect AI governance discussions with international efforts, such as the International Independent Scientific Panel, the policy dialogue on AI within the United Nations, and future initiatives like the upcoming International AI Standard Summit in Seoul."*



Figure 7: Amandeep Singh Gill, Under-Secretary-General and Special Envoy for Digital and Emerging Technologies, Office for Digital and Emerging Technologies, United Nations

Sung Hwan Cho, President, ISO - *"Inclusion and diversity are at the core of international standards' goal. We have to ensure no one is left behind. AI standards should cover technology, but also how AI benefits humanity."*



Figure 8: Sung Hwan Cho, President, ISO

Seizo Onoe, Director of the ITU Telecommunication Standardization Bureau - *“The World Standards Cooperation partnership is key to comprehensive standards development for AI. AI has created even stronger connections among standards bodies, and AI is evolving very fast. We want to ensure that our standards keep pace with this evolution.”*

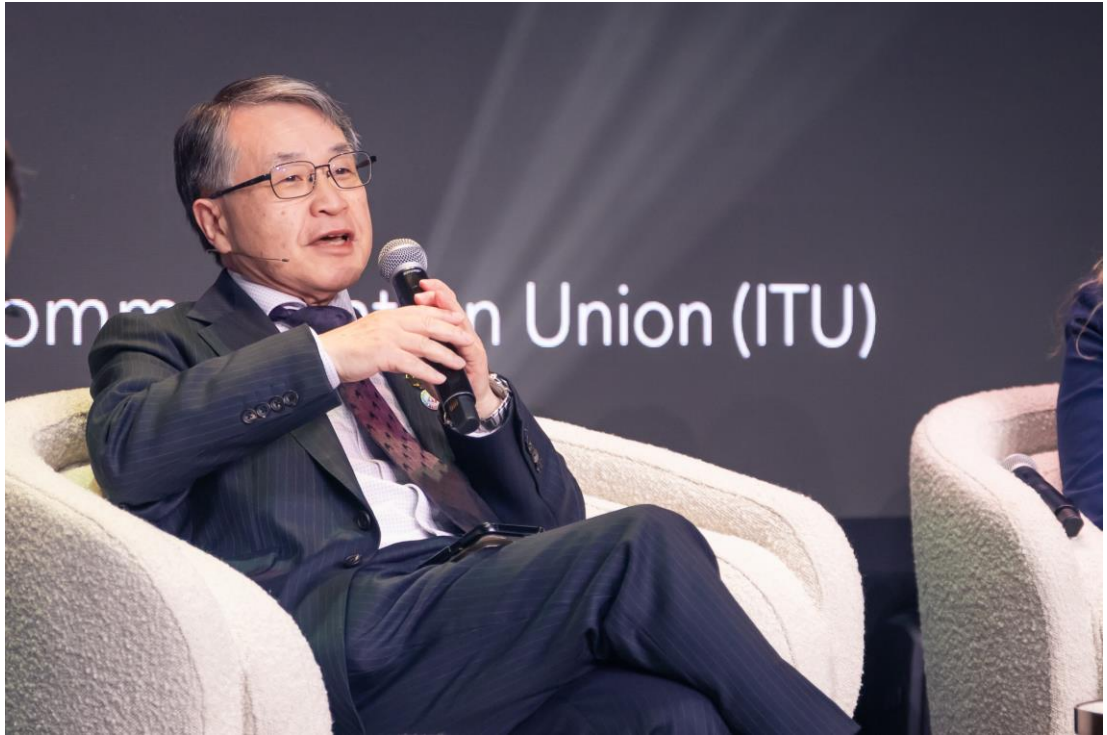


Figure 9: Seizo Onoe, Director of the Telecommunication Standardization Bureau, ITU

Philippe Metzger, CEO and Secretary General, IEC – *“Standards development organizations ultimately translate principles into practical implementation. Standards bodies need to collaborate in a structured way to ensure the clarity and coherence essential to global impact.”*



Figure 10: Philippe Metzger, Secretary-General & CEO, International Electrotechnical Commission (IEC)

Kathleen Kramer, President and CEO, IEEE - *"Achieving impactful standards requires intentional, systematic, and open collaboration that involves all stakeholders, ensuring the outcomes are technically sound, widely accepted, and adaptable to emerging technologies over time."*



Figure 11: Kathleen A. Kramer, President & CEO, IEEE

Paul Gaskell, Deputy Director of Digital Trade, Internet Governance & Digital Standards at the Department for Science, Innovation & Technology, UK - *"The UK advocates for global coordination among SDOs to avoid duplication and fragmentation. Initiatives like this, the UK AI Standards Hub and a recent global summit on AI standards in London demonstrate the UK's commitment to collaboration with organizations such as the ITU, ISO, and IEC."*



Figure 12: Paul Gaskell, Deputy Director of Digital Trade, Internet Governance & Digital Standards at the Department of Science, Innovation & Technology, UK

Panelists emphasized the need for faster, more responsive standards development for AI by fostering tighter collaboration among global, national, and regional standards bodies, leveraging their unique strengths to meet shifting demands effectively.

The AI Standards Exchange Database provides a single, transparent, and traceable access point for AI-related standards, benefiting policymakers, regulators, companies, and developers by offering up-to-date insights on state-of-the-art technology.

4 High-level roundtable - From principles to practice: how AI standards enable effective governance

As AI continues to mature and scale globally, the gap between governance frameworks and technical standardization efforts has become increasingly evident. While both are critical to ensuring trust, accountability, and scalability, they too often evolve in parallel rather than in partnership. This high-level roundtable convened 100 leaders from governments, standards bodies, and companies to explore how to bridge this divide and identify the concrete steps needed to align governance and standardization in shaping the future of AI.

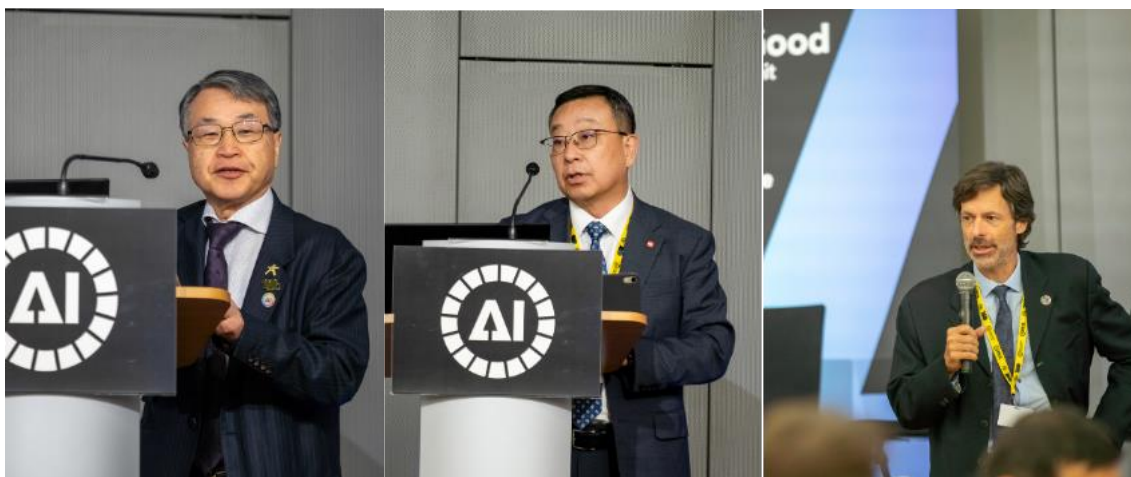


Figure 13: Left to right: Seizo Onoe, Director of the Telecommunication Standardization Bureau (TSB), International Telecommunication Union (ITU); Sung Hwan Cho, President, International Organization for Standardization (ISO); Philippe Metzger, Secretary-General & CEO, International Electrotechnical Commission (IEC)

Following remarks from representatives of ITU, ISO and IEC, an "Insight Sparks" segment featured talks on AI standards, governance, and achieving synergy between the two from Hatem Dowidar, Group CEO at Etisalat; Sasha Rubel, Head of Public Policy for Generative AI at Amazon Web Services (AWS); and Xiao Ran, Vice President at Huawei Technologies.

A common message from the high-level speakers was that trust is the backbone of AI adoption. For people to rely on AI in sectors like healthcare, finance, mobility, and government, technical AI standards are needed for trust, transparency, accountability, upholding human rights, and helping ensure no one is excluded.

Roundtable discussions following the Insight Spark segment addressed five questions.

Q1. How can standards bridge AI principles with regulatory and legislative compliance in practice?

- Standards can help align AI principles with regulatory compliance, said participants, noting that a design framework rooted in principles like transparency and trustworthiness could guide regulation, using examples like government action on deepfakes.
- Participants highlighted the need for clarity in AI-related definitions across legislative frameworks (e.g. EU AI Act, EU General Data Protection Regulation (GDPR), U.S. State laws),
- Participants emphasized that standards can be forward-looking, inclusive of diverse stakeholders, and complementary to national legislation, particularly in sectors like education, energy, and health, while also helping harmonize approaches to issues like synthetic media across markets.
- Participants emphasized the need for inclusive frameworks where regulators and legislators can participate early in the development of AI standards to bridge the gap between technology advancement and regulation, with a focus on skills development as a foundational element.

- Participants highlighted the challenges of aligning legislative aspirations, such as those in the GDPR, with practical business practices, using cookie consents as an example, and noted that this misalignment becomes even more complex with AI due to its diverse applications and the need to balance ethical principles with technological advancement.
- Participants also noted the gap between ambitious AI principles (e.g., explainability) and the current maturity of technologies needed to achieve them, questioning how standards can effectively bridge this gap.



Figure 14: Bilel Jamoussi, Deputy to the Director, Telecommunication Standardization Bureau (TSB), International Telecommunication Union (ITU)

Q2. What AI risks must be prioritized for global standardization today?

- Participants discussed prioritizing risks for AI standardization, highlighting three key areas: information authenticity; data security throughout data lifecycles; and AI reliability to address trust, accountability, and hallucination issues. Participants also noted that standards often emerge naturally from industry needs alongside risk-focused efforts.
- Participants identified key risks for AI standardization, including data quality and accuracy (especially in autonomous driving) and misinformation/deepfakes (validated as a major concern by AI tools), and noted that while standards can provide definitions and frameworks, they cannot address the root sources of misinformation directly.
- Participants highlighted an urgent need for global standardization at the intersection of AI and cybersecurity, emphasizing risks related to intellectual property, critical infrastructure (e.g. telecommunications), cybercrime, autonomous driving, and social

networks, with standards serving as a means for secure, compliant operations and governance.



Figure 15: Vijay Mauree, Programme Coordinator, Telecommunication Standardization Bureau (TSB), International Telecommunication Union (ITU)

Q3. How can global interoperability for AI systems work without undermining national regulatory sovereignty?

- Participants emphasized the balance between AI sovereignty and global collaboration, highlighting the value of interoperability through voluntary frameworks for responsible AI (e.g. for safety, accountability, and sustainability) and technical standards, while respecting national laws and addressing data privacy through potential data interoperability frameworks.
- Participants stressed the importance of multistakeholder collaboration for interoperability, emphasizing the need for consensus-driven, primarily technical standards while considering governments' policy concerns and frameworks like the Organization for Economic Cooperation and Development (OECD) principles and the Council of Europe Convention.
- Participants agreed on the need for global AI interoperability and defining AI agent interactions, security, and trust properties, while recognizing that issues like cross-border cloud operations already challenge national regulatory sovereignty. Participants also highlighted a new concern around AI agents autonomously communicating in ways that could pose risks to data integrity or cybersecurity, for example.
- Participants advocated for transnational governance rooted in existing global frameworks (e.g. Universal Declaration of Human Rights and Convention on the Rights of the Child), emphasizing the importance of a human-centred approach, broad multistakeholder collaboration, and consensus on shared principles to guide AI development and lifecycle governance.



Figure 16: Sasha Rubel, Head of Public Policy for Generative AI, Amazon Web Services (AWS) High-level roundtable - From principles to practice: How AI standards enable effective governance

Q4. What role can standards play in ensuring AI supports inclusive, sustainable development?

- Participants advocated for global AI standards as equity enablers to promote algorithmic fairness, explainability, and non-discrimination; emphasized that standards can serve as tools for inclusion through participation in standards development, particularly for the Global South, while supporting interoperability and contextual relevance; highlighted standards' value as drivers of responsible innovation, helping to ensure sustainable scaling by managing risks like privacy and security; and positioned standards as a shared language for cross-border, cross-sector coordination, stressing the need for inclusive, co-created frameworks aligned with sustainable development goals and ethical foresight.
- To help ensure AI inclusivity, said participants, open-source models could be paired with bias and sustainability metrics to account for differences in country-specific training data, supported by regulatory sandboxes and active collaboration across academia, governments, and the broader ecosystem.
- Participants highlighted the potential future importance of large-scale AI models, suggesting that conflicts might target AI systems rather than traditional infrastructure, underscoring the need for standards that balance global coordination with national sovereignty.
- Participants highlighted the need for standards to support data governance and help ensure data quality (e.g. using sensors for measurement), improve AI model training

efficiency, optimize resource distribution, and enable the reuse of existing datasets. They also emphasized the need for creating standards to measure AI energy consumption and defining interfaces to help small companies, researchers, and industry partners integrate with large AI platforms effectively.



Figure 17: Xiao Ran, Vice President, Huawei Technologies Co., Ltd.

Q5. What multistakeholder model is needed to accelerate trusted AI standards and what are the roles of different stakeholders in the process?

Participants discussed the importance of involving diverse stakeholders (e.g. government, industry, academia, civil society, standards bodies, and users) in standardization processes, while also acknowledging challenges related to inclusivity, conflicting interests and engagement in complex technical discussions. Participants also emphasized the need for proactive regulation, clear role definitions, and practical tools (e.g. checklists) to guide AI producers in preventing harm and ensuring accountability throughout the product lifecycle.



Figure 18: Hatem Dowidar, Group CEO, e&

5 High level panels and keynotes

5.1 Transforming telecoms with AI and machine learning

As a tech-native industry, the telecommunications sector has embraced machine learning (ML) and predictive AI to optimize network operations for over a decade. Today, with the rapid advancement of generative AI, telecom operators are entering a new phase – reimagining traditional business and operating models. AI is enhancing network performance, reducing energy consumption, and driving efficiencies at scale. From intelligent automation that lowers cost-to-serve to generative AI-powered personalization in sales and marketing, AI is unlocking innovative services, new business models, and responsive, dynamic networks. There is also growing momentum behind telecom-led AI offerings such as AI infrastructure provision and AI-as-a-service (AIaaS), while natural language capabilities are reshaping customer service through digital assistants and chatbots. In parallel, telecom operators are increasingly positioning themselves as trusted partners in the AI age.

John Omo, Secretary-General of the African Telecommunications Union (ATU); Hatem Dowidar, Group CEO of Etisalat; Chih-Lin I, Chief Mobile Scientist at China Mobile; and Pamela Snively, Chief Data and Trust Officer at TELUS, discussed real-world use cases already delivering impact and highlighted the growing importance of international standards and collaboration in shaping the future of AI-powered telecommunications.

Some of the key points discussed and raised by panelists are summarised below:

a. Rapid pace of AI evolution

- Over the past decade, leveraging wireless big data and evolving AI technologies has significantly improved resource efficiency and performance in telecommunications. For example, advanced AI now enables over 1 billion kWh in annual energy savings by optimizing multi-generation (2G/3G/4G) networks across over 5 million base stations, demonstrating immense cost-saving opportunities.
- AI-powered tools, such as speech recognition and automatic translation, are creating innovative services that improve daily lives. For instance, in China, AI supports communication across 56 ethnic groups, eight dialect families, and 30 accent groups, facilitating seamless interaction and dialogue translation. Additionally, as part of progress towards 6G, networks should both leverage AI for optimization and enable creative, socially beneficial AI applications, effectively offering "AI as a service."
- The success of the mobile industry, from 1G to 5G and now preparing for 6G, is attributed to globally unified standards that help ensure economies of scale, global interoperability, and societal transformation. Similarly, in the rapidly evolving AI landscape, global standards are essential to establish trust, reliability, and confidence in the technology.
- Unlike the traditional 10-year cycle in mobile communication standards, AI's fast-paced evolution demands a more dynamic, continuous approach to global standardization. This can help ensure interoperability, scalability, and trust while adapting to AI's rapidly shifting advancements.

b. AI and ML applications in telecoms

- Telecommunications companies – with their expertise in deploying large-scale technology, managing data, and safeguarding privacy – are well positioned to drive discussions on AI standards to meet address societal and community needs effectively.
- AI, particularly generative AI, is revolutionizing telecommunications by improving customer care experiences (e.g. moving beyond outdated chatbots) and optimizing operations such as self-healing networks, network capacity management, and energy consumption.
- Modern telecommunications infrastructure, with advancements like low-latency networks and edge computing, supports AI by enabling data processing in the cloud and facilitating previously impossible applications.
- AI enables real-time network reconfiguration and load balancing, addressing equipment failures seamlessly without customer impact. This improves user experience, prevents revenue loss, and eliminates the need for manual interventions that used to take hours.
- In five years, advancements in agent AI within networks and digital infrastructure could drive greater efficiency, sustainability, and equitable access to knowledge and

skills, supported by collaborative efforts from organizations like ITU and other United Nations organizations.

- AI enhances telecom operations by automating tasks like bill reviews, reducing errors, and optimizing resources. It also drives demand for 5G's low latency, powering applications like autonomous vehicles, drones, and private networks that add significant value to the ecosystem.

c. Importance of AI standards

- The growing AI digital divide, driven by data sovereignty and unequal access to AI capabilities, risks widening disparities between operators in advanced markets and those in regions like Africa, underscoring the need for standards to support shared global progress.
- Standardization plays a critical role in enabling interoperability, energy efficiency, security, trustworthiness, and cooperation across various networks and systems.
- Collaboration among standards bodies is essential to developing standards that enhance scalability, interoperability, and the security of networks, addressing the evolving needs of industry and customers.
- Open standardization processes enable broad participation and help ensure localized relevance and inclusivity, while more exclusive processes may limit value and alignment with diverse global needs.



Figure 19: Left to right: Ebtesam Almazrouei, Executive Director of the Office of AI and Advanced Technology at the Department of Finance, CEO and Founder of AIE3, Chairperson of UN AI for Good Impact Initiative. John Omo, Secretary-General, African Telecommunications Union (ATU). Chih-Lin I, China Mobile Research Institute. Pamela Snively, Chief Data & Trust Officer, TELUS. Hatem Dowidar, Group CEO, e&.

5.2 AI Standards: Building trust and supporting innovation in a networked world

Cameron F. Kerry, *Ann R. and Andrew H. Tisch Distinguished Visiting Fellow, The Brookings Institution*, highlighted the following in his keynote:

- The unpredictability of AI advancements, mentioned by the UK Safety Report and other experts, necessitates a flexible and adaptive approach.
- Standards were identified as the number one priority for global AI collaboration, ranking above ethical alignment, safety, and research and development (R&D).
- Standards bodies like ITU, ISO, and IEC play a critical role in creating frameworks for transparency, accountability, interoperability, and risk mitigation.
- Standards bodies provide tools for measurement and assessment, helping to operationalize AI policies and related best practices.
- The distributed and multistakeholder approach used in internet governance (e.g. Internet Corporation for Assigned Names and Numbers (ICANN)) was highlighted as a model for AI governance.
- When it comes to AI safety and governance, some of the important work to be considered – which show how principles evolve into frameworks, standards, or binding rules – include the following, among others:
 - The OECD AI Principles (2019), which laid the foundation for global AI ethics.
 - The Hiroshima AI Principles, initiated by Japan during the G7, which emphasize transparency and accountability in foundational models.
 - The Global Partnership on AI (GPAI) and the Hiroshima Friends Group, which now includes 56 governments, demonstrating the global uptake of AI governance efforts.
- AI governance should mirror the distributed, resilient, and adaptive nature of neural networks and internet infrastructure models.
- A multi-layered, multi-stakeholder, and multifaceted approach is required to address the complexities of AI across different contexts and regions. Broad participation from governments, civil society, and underrepresented stakeholders is essential for inclusive governance.
- Standards bodies should prioritize transparency and a diversity of voices in their processes to address socio-technical challenges effectively.
- The development of measurement and assessment tools is critical for AI governance, enabling continuous monitoring and adaptation throughout AI systems' lifecycles.
- Standards provide a foundation for interoperability, facilitating collaboration across differing legal and regulatory systems.
- ITU's diverse global membership and expertise in capacity building positions ITU as a key player in advancing global AI standards.



Figure 20: Cameron F. Kerry, Ann R. and Andrew H. Tisch Distinguished Visiting Fellow, The Brookings Institution

5.3 Navigating tomorrow: Education, skills, and standards in an AI-driven workplace

Beena Ammanath, Author “Trustworthy AI” & “Zero Latency Leadership” and Global Deloitte AI Institute leader, gave a keynote talk highlighting the importance of lifelong learning, robust standards, and human-AI collaboration to build a more inclusive, innovative future of work, with key points including:

- Over the past few years, AI has made remarkable advancements, transitioning from niche use cases to widespread applications in diverse fields such as healthcare, customer service, and drug discovery. AI is now an integral part of everyday life, offering transformative potential for good.
- The ability of an organization’s workforce to understand and effectively use AI will be a key competitive advantage in the future. This requires a focus on workforce education, skills development, and standards to ensure employees are equipped to work alongside AI tools.
- The rapid evolution of AI demands a shift from one-time education to lifelong, continuous learning. Employees must have a foundational understanding of AI, paired with ongoing training to stay updated with new tools and technologies.
- Personalized AI-enabled training methods, such as AI coaching, can make learning more engaging and tailored to individual needs and roles.

- Employees trained in AI tools can improve efficiency and create opportunities for innovation. New roles, such as data scientists and AI coaches, will emerge, but this requires individuals to embrace AI proactively and integrate it into their work.
- The "three lanes of AI development":
 - AI research: Core advancements in AI technologies, such as LLMs.
 - Applied AI: Real-world implementation of AI in industries such as banking, healthcare, and supply chains, even as the technology continues to mature.
 - Standards and guardrails: Developing regulations, safety standards, and ethical frameworks to ensure AI is used responsibly and effectively.
- Trust in AI adoption depends on clear standards, safety regulations and risk management practices. Standards can act as the "seat belts and speed limits" for AI, helping ensuring responsible development and use of the technology.
- Employees at all levels should be trained to recognize and mitigate risks associated with AI tools. This includes understanding regulatory compliance and using AI responsibly to avoid unintended consequences, such as mishandling proprietary data.
- Organizations should provide safe environments for employees to experiment with AI tools. AI sandboxes allow employees to learn, test, and familiarize themselves with new technologies without risking compliance or data security.
- A Call to Action was made to leaders to prioritize AI education, create pathways for continuous learning, and foster a culture of responsible AI use. Employees should proactively learn about AI tools relevant to their roles to remain competitive in the evolving workplace



Figure 21: Beena Ammanath, Author “Trustworthy AI” & “Zero Latency Leadership” and Global Deloitte AI Institute leader, Deloitte

5.4 Computing and AI: Endless frontier and exploration

Wang Jian, Director at Zhejiang Lab gave a keynote talk highlighting the following areas where computing and AI will play a key transformative role:

- Relationship between computing and AI:
 - Computing and AI are deeply interconnected, akin to two sides of the same coin. Alan Turing's early work emphasized computing as a tool to extend human intelligence, much like pen and paper. Today, AI is amplifying human creativity and problem-solving abilities.
 - Computing is not just a subset of computer science but a fundamental discipline parallel to fields like physics and life sciences, forming the foundation upon which AI is built.
- AI as a transformative tool for science:
 - AI is revolutionizing scientific research by enabling scientists to achieve breakthroughs that would have been impossible otherwise. Examples include projects like GeoGPT, an AI system designed to assist geoscientists in data sharing, research collaboration, and discovery.
 - The Deep Time Digital Earth Initiative and related projects demonstrate how AI can provide innovative tools for researchers, with applications in geoscience, classification, and open science.
- AI and space exploration:
 - The speaker envisions a future where AI and computing are integral to space exploration. Initiatives like the Three-Body Computing Constellation aim to process multi-source data directly in space, reducing the need to send data back to Earth.
 - Proposed computing satellites would complement existing satellite types (communication, navigation, and observation) to perform advanced in-orbit processing, enabling real-time collaboration for phenomena like gamma-ray bursts and solar observations.
- AI as a fundamental scientific tool
 - AI is becoming as essential to science as mathematics, applicable across various disciplines. The development of domain-specific foundation models tailored for scientific data (beyond text) will further enhance AI's role in advancing research and innovation.
 - Shared infrastructures, such as the open platform 02x.org, aim to democratize access to AI capabilities and foster global collaboration among scientists.
- Governance and safety in AI

- Governance is critical for ensuring AI and computing technologies are used responsibly. The GeoGPT governance committee is a notable example of how structured oversight can address safety, ethical, and intellectual property (IP) concerns in scientific applications of AI.
- Open science practices, such as those demonstrated in GeoGPT projects, highlight the importance of transparency and collaboration in advancing technology for the benefit of humanity.
- Global challenges and collaboration:
 - AI and computing can help address pressing global challenges, such as climate change and disasters stemming from natural hazards, by enabling deeper understanding and innovative solutions. The vision of Earth intelligence is an example of how AI can improve planetary observation and problem-solving.
 - Collaboration across organizations, countries, and disciplines is essential for tackling complex challenges. Initiatives like the Three-Body Computing Constellation emphasize the importance of collective effort to achieve ambitious goals.
- The future of AI and computing:

AI and computing will be indispensable companions for humanity's journey to Mars and beyond. AI's capabilities will continue to redefine what is possible, both on Earth and in space, driving innovation and exploration.



Figure 22: Wang Jian, Director, Founder of Alibaba Cloud, Zhejiang Lab

5.5 AI and multimedia authenticity

Alessandra Sala, Senior Director of AI Data Science at Shutterstock and Chair of the AI and Multimedia Authenticity Standards Collaboration (AMAS) presented the key achievements of AMAS and highlighted the following in her intervention:

- a) Generative AI has transformative potential but also introduces risks like misinformation and disinformation, identified by the World Economic Forum and the UN as top global risks. Examples of risks include:
 - i. Deepfakes used for gender violence: Non-consensual content targeting women and activists in vulnerable regions.
 - ii. Exploitation of children: A 380% rise in child abuse content reported by the UK Internet Watch Foundation.
 - iii. Political manipulation: Deepfakes used to polarize societies.
- b) These risks threaten societal cohesion, safety, and trust, highlighting the need for multilateral collaboration.
- c) AMAS was established under the World Standards Cooperation (ITU, ISO and IEC) and announced at the AI for Good Global Summit in 2024. The main objective is to coordinate areas where technical standards and policies are needed to address risks related to deepfakes and create awareness on these issues to bring them to the attention of policymakers and regulators.



Figure 23: Alessandra Sala, Global President of Women in AI, Chair of AI and Multimedia Standards Collaboration, and Sr. Director of Artificial Intelligence and Data Science, Shutterstock.

- d) Key achievements and deliverables launched:
 - i. [Technical Paper on AI and Multimedia Authenticity Standards](#): A landscape analysis of over 35 existing standards, which are classified into six categories:

content provenance, trust and authenticity, asset identifiers, rights declarations, and watermarking.

- ii. [Policy Paper: Building trust in Multimedia Authenticity through International Standards](#): A practical tool for policymakers and regulators, developed after studying AI regulation, policy, and practices for multimedia authenticity across 60 countries. It includes recommendations and a checklist for regulators and policymakers to address gaps and challenges in regulating AI and multimedia authenticity.
- e) Future plans:
- i. Addressing gaps in existing standards and coordinating efforts across multiple standard-setting organizations.
 - ii. Upcoming events include the International AI Standards Summit in the Republic of Korea, 2-3 December 2025 where further updates, discussions, and coordination efforts will be shared.



Figure 24: Left to right: Leonard Rosenthol, Senior Principal Architect, Adobe and Vice Chair of AMAS; Touradj Ebrahimi, Professor, Ecole Polytechnique Fédérale de Lausanne, Convenor of JPEG Group, and Vice Chair of AMAS; Alessandra Sala, Global President of Women in AI, Chair of AMAS, Sr. Director of Artificial Intelligence and Data Science, Shutterstock; Cindy Parokkil, AI Policy Lead & Programme Manager, ISO Central Secretariat and AMAS Vice Chair; Mike Mullane, Deputy Director of Communications, International Electrotechnical Commission (IEC) and AMAS Vice Chair, and Vijay Mauree, Programme Coordinator, ITU and AMAS Secretariat

5.6 Humans and AI: The Journey

Manuela Veloso, Head of AI Research at JPMorgan Chase, explored the evolving interaction between humans and AI from her experience in AI in robotics and AI in finance, highlighted the following key aspects in her keynote talk:

- Integration of AI capabilities:
 - AI integrates three key modules: perception, cognition, and action, aiming to mirror the way humans operate intelligently.
 - Perception involves understanding the environment (e.g. sensors and language processing). Cognition includes decision-making and reasoning. Action involves executing decisions, such as moving or communicating.
- Autonomous AI agents:
 - Examples like robots playing soccer demonstrate AI systems operating autonomously, making real-time decisions without human input.
 - These agents process sensory data, strategize, and act to achieve goals without pre-programmed moves.
- AI with limitations:
 - Some AI systems, like Carnegie Mellon's cobot robots, exhibit limitations (e.g. inability to press elevator buttons).
 - To overcome such constraints, these systems ask humans for help.
 - This symbiosis between humans and AI highlights the importance of collaboration, where humans assist AI in completing tasks.
- Learning from human Interaction:
 - Collaborative AI systems learn from human assistance. For example, a robot learns the location of objects or tasks based on human input, improving over time.
 - This concept applies to digital agents as well. For instance, a "Docubot" creates PowerPoint slides based on user requests and learns from user feedback to refine its outputs.
- Digital agents:
 - Digital agents integrate perception, cognition, and action in a virtual context.
 - They perform tasks such as reading documents, reasoning with data, and generating actions (e.g. creating reports or sending emails).
 - These agents ask for help when they encounter ambiguous instructions, learn from human responses, and adapt to future tasks.
- Generative AI as a mediator:
 - Generative AI plays a critical role in translating between human language and AI systems' internal logic.
 - Generative AI enables seamless communication between humans and AI agents, fostering collaboration and improving task execution.

- AI development philosophy:
 - AI systems should acknowledge their limitations and ask for help when uncertain, rather than attempting to solve everything autonomously and risking errors (e.g. hallucinations in ChatGPT).
 - Continuous learning and adaptation are crucial for building robust AI systems. AI can improve by observing how humans solve problems and retaining this knowledge for future use.
- The future of human-AI collaboration:
 - The vision for the future includes a symbiotic ecosystem of humans and multiple AI agents.
 - AI agents will interact with each other (e.g. weather agents and social media agents) and with humans to solve problems collaboratively.
 - Humans will play a key role in guiding and teaching AI systems while leveraging their capabilities to automate tasks, solve complex problems, and scale solutions.
- The ultimate goal is to build AI agents that can learn, adapt, and collaborate with humans, creating a harmonious coexistence. This requires embracing the fact that AI does not know everything upfront but improves through continuous interaction and feedback.
- AI and humans will increasingly coexist in a mutually beneficial relationship, with AI systems learning from humans and humans leveraging AI's capabilities to enhance productivity and innovation.



Figure 25: Manuela Veloso, Head, AI Research, JPMorgan Chase

5.7 Donate your brainwaves for social impact

This session featured two keynote speakers, Rodrigo Hübner Mendes, CEO and Founder of Institute Rodrigo Mendes, and Olivier Ouillier, Co-founder and CEO of Inclusive Brains, Visiting Professor at Mohamed Bin Zayed University of Artificial Intelligence (MBZUAI), and Chairman of the AI Institute at Biotech Dental Group.

By becoming the first person ever to mind-control a real Formula One car on a real racetrack using a non-invasive Brain-Computer Interface (BCI), Rodrigo Hübner Mendes has inspired millions around the world and shown that AI-powered assistive technologies can be life-changing. AI converted his brain waves into commands, allowing him to accelerate, turn, and control the car on a racetrack. This groundbreaking achievement captured attention, including from Formula One legend Lewis Hamilton, who accepted Rodrigo's challenge to race using the same mind-control technology.

In partnership with Allianz Trade, Inclusive Brains, Olivier Ouillier and his team developed Prometheus BCI, a multimodal AI interface that enabled three non-invasive BCI world premieres in the past 12 months:

- The first mind-controlled torch relay during the Paris 2024 Olympics.
- The first tweet exchange (without any vocal command or physical interaction) with French President Emmanuel Macron, live from the stage of the 2024 AI for Good Summit.
- The first-ever parliamentary amendment finalized and sent "hands off" during the AI Action Summit in Paris to the President of the Assemblée Nationale Yaël Braun-Pivet.

Traditionally, humans adapt to machines, but the vision is to create systems that adapt to users by understanding their unique attributes. Olivier Ouillier combines data science, neuroscience, and behavioural science to build adaptive systems using brain data, facial expressions, eye tracking, heartbeat, and movement. Olivier and his team adapted their BCI technology to help individuals with disabilities. For example, Natalie, a person with cognitive impairments, became the first Olympic torchbearer to use mind-control technology, combining brain waves, facial expressions, and heartbeat to control an exoskeleton arm. The technology was later improved and used by others, who showcased advanced control and interaction with crowds during Olympic events.

Practical applications extend beyond disabilities, for example:

- Surgeons tracking stress and cognitive load in lengthy procedures to improve performance and safety.
- Creating stress-monitoring systems for dental surgeons and patients, showcasing its broad applicability.

The team has made their BCI code open source to encourage widespread innovation and collaboration. This move aims to scale the technology and make it affordable, enabling broader access and benefiting more people.

The Brain Wave Donation Campaign objective is to advance inclusive AI, encouraging people to donate their brain waves for social good. The goal is to build a large open-source database of brain waves to train AI systems for assistive technologies. Donors retain ownership of their data, ensuring privacy and control over how their contributions are used.

The ultimate mission is to create inclusive technologies that benefit everyone, starting with those who have disabilities. By scaling and refining the technology, the team aims to provide solutions that promote education, employment, and accessibility for people with impairments. The benchmark is to create technology akin to the remote control, which started as an assistive device but became universally adopted.

A Call to Action was made inviting researchers, innovators, and the public to contribute to the brain wave database and collaborate on developing inclusive AI.

In conclusion, adaptive AI and brain-computer interfaces hold transformative potential to make technology more inclusive and accessible, benefiting both individuals with disabilities and society at large. The focus is on collaboration, open innovation, and ethical data use to drive meaningful social impact.



Figure 26: Rodrigo Hübner Mendes, CEO & Founder, Institute Rodrigo Mendes & Olivier Ouillier, Co-Founder & CEO, Inclusive Brains; Visiting Professor, Mohamed Bin Zayed University of Artificial Intelligence (MBZUAI); Chairman, AI Institute, Biotech Dental Group.

5.8 AI for food systems

This high-level session between International Telecommunication Union (ITU), the Food and Agriculture Organization of the United Nations (FAO), the World Food Programme (WFP), and the International Fund for Agricultural Development (IFAD) discussed how to leverage the transformative power of AI and digital technologies across the food value chain to enhance productivity, strengthen food security, and build resilient food systems building on the thematic workshop held earlier. The summit also hosted the announcement of the launch of a new Global Initiative on AI for Food Systems, a strategic partnership between ITU, FAO, WFP, and IFAD.

ITU's partners were represented by Dejan Jakovljevic, Director of the Digital FAO and Agro-Informatics Division; Pieter Boogaard, Managing Director of Technical Delivery at IFAD; Magan Naidoo, Chief Data Officer at WFP; and Sebastian Bosse, Head of Interactive and Cognitive Systems at Fraunhofer Heinrich Hertz Institute, and Chair of the ITU/FAO Focus Group on AI and IoT for Digital Agriculture – the predecessor of the new Global Initiative.

The main points shared during the session are summarised below:

- The global food system faces unprecedented challenges, including climate shocks, demographic pressures, and supply disruptions, necessitating innovative solutions.
- The new Global Initiative aims to leverage AI to provide practical, scalable solutions aimed at addressing humanity's urgent food security needs through sustainable and inclusive approaches.
- The initiative emphasizes global collaboration to create interoperable, secure, and adaptable frameworks for digital food systems.
- The initiative aims to empower governments and innovators to create resilient food systems, ensuring no one, particularly small farmers and vulnerable communities, is left behind in the digital age.
- The session highlighted the importance of ethical, equitable, and responsible AI development to transform food systems and improve the lives of smallholder farmers, who produce a third of the world's food but often lack access to resources and support. The Global Initiative aims to address these challenges and foster shared prosperity, equity, and sustainability.
- The session highlighted the critical role of AI in addressing global food insecurity, improving food systems, solving logistical challenges, and building resilience in agriculture.
- The session underscored the urgency of transforming agri-food systems to address the needs of the 700 million people facing food insecurity, identifying AI as a critical enabler and accelerator for driving these changes effectively and responsibly.
- The session emphasized the importance of collaboration on horizontal issues such as standardization, framework structuring, and reference architectures for responsible AI

use, highlighting the potential of the Global Initiative to elevate joint efforts into actionable solutions for food system transformation.

- The Global Initiative builds on the work of the ITU/FAO Focus Group on AI and IoT for Digital Agriculture, which has made significant progress in advancing digital agriculture through open interoperability and discrimination-free practices.



Figure 27: Left to right: Pieterneel Boogaard, Managing Director, Office of Technical Delivery, United Nations-International Fund for Agriculture Development (IFAD); Magan Naidoo, Chief Data Officer, United Nations World Food Programme (WFP); Seizo Onoe, Director of the Telecommunication Standardization Bureau (TSB), International Telecommunication Union (ITU); Dejan Jakovljevic, Chief Information Officer, Director of Digital FAO and Agro-Informatics Division, FAO; Sebastian Bosse, Head of the Interactive & Cognitive Systems Group, Fraunhofer Heinrich Hertz Institute (HHI)

5.9 Future of smart mobility

Talks on the future of smart mobility, focusing on the collaborative efforts between ITU and the UN Economic Commission for Europe (UNECE) to drive innovation in the sector, considered the key takeaways from ITU-UNECE Future Networked Car Symposium, shedding light on how emerging technologies like AI are transforming transportation.

The session welcome Francois E. Guichard, UNECE Focal Point for Intelligent Transport Systems and Automated Driving; Helen Pan, General Manager of Baidu Apollo; and Vincent Vanhoucke, Distinguished Engineer at Waymo.

A glimpse into the future of smart mobility was analysed with Baidu and Waymo providing insights on key areas of focus for the future of smart mobility.

Some of the key highlights of the session were:

- Baidu's autonomous driving technology has achieved significant milestones, including over 170 million km of public road testing, a 10X safety improvement over human drivers, and 14X fewer insurance claims with over 1,000 driverless vehicles globally, while prioritizing safety, trust, and collaboration with regulators for cautious expansion.
- Beyond safety, Baidu focuses on sustainability, accessibility, and inclusivity by integrating its robotaxi service with public transportation, addressing underserved mobility needs (e.g. rainy days and rush hours) and providing tailored solutions for women, senior citizens, and vision-impaired users through features like touchless entry, voice commands, and service dog accommodations.
- Autonomous driving technology uniquely requires real-time problem-solving, high accuracy, and reliability, leveraging advanced AI algorithms and reasoning models. As early adopters of AI advancements, the industry continues to integrate new innovations to commercialize safer and more effective solutions.
- Waymo reported that autonomous driving technology has received overwhelmingly positive public feedback, becoming a popular and trusted option across demographics, with benefits like increased safety and convenience for women, parents, and teens, as well as consistent and independent transportation experiences.
- With over 50 million miles driven, data shows significant safety improvements, including a 92% reduction in pedestrian injuries and an 82% reduction in injuries to cyclists and motorcyclists, highlighting the potential to save lives on a large scale as the service expands.
- Waymo's successful launch of fully autonomous services in five cities, with increasingly faster expansions, highlights progress in overcoming logistical and community trust challenges as the focus shifts to scaling operations, building infrastructure, and collaborating with local authorities for broader U.S. expansion.
- Both Waymo and Baidu consider autonomous cars are valuable to safety and consistency for pedestrians and cyclists, while also reducing traffic issues like congestion, leading to smoother and more efficient urban mobility systems.
- Waymo's autonomous vehicles have the potential to transform accessibility for individuals with disabilities or temporary mobility challenges, offering newfound independence, privacy, and autonomy to communities traditionally excluded from driving.

- The evolution of autonomous driving highlights its shift from a robotics challenge to an AI problem, requiring human-like intelligence and intuition to navigate complex, context-driven "visual conversations" on the road, leveraging advanced AI tools like multimodal models, world models, and reinforcement learning from human feedback.
- The rapid evolution of AI, including advancements in LLMs, multimodal models, and world models, promises transformational impacts on autonomous driving, with constant reinvention driving progress through successive generations of machine learning technologies over the past 15 years.



Figure 28: Left to right: Francois E. Guichard, Focal Point, Intelligent Transport Systems and Automated Driving, United Nations Economic Commission for Europe (UNECE); Helen Pan, General Manager, Baidu Apollo; Vincent Vanhoucke, Distinguished Engineer, Waymo

5.10 AI and energy

Gitta Kutyniok, Bavarian AI Chair for Mathematical Foundations of Artificial Intelligence, at Ludwig-Maximilians Universität München, and Qi Shuguang, Vice Deputy Engineer at the China Academy of Information and Communications Technology (CAICT); shared the key outcomes of the summit's workshop on "Navigating the Intersect of AI, Environment and Energy for a Sustainable Future," highlighting the progress in the area of energy efficiency for AI and upcoming challenges in the future where standards would be needed.

AI can play a significant role in climate monitoring and reducing environmental impacts, including energy and water consumption and GHG emissions. It can also help improve

sectoral systems such as power grids, agriculture, waste management, biodiversity conservation, and transport and mobility. For example, in the manufacturing and energy sectors, AI contributes to energy savings and supports green transitions through device control, process optimization, recycling, IoT integration, and deep learning technologies.

The key points discussed are summarised below:

- a. Challenges for energy efficiency for AI
 - The environmental impacts of AI – including significant energy consumption, CO₂ emissions, and health risks from PM 2.5 near hyperscale data centres – is growing exponentially with advancements like generative AI and cloud computing, which rely on massive data centres.
 - There is a pressing need for new metrics and independent scientific assessments to evaluate AI's environmental footprint.
 - The industry faces significant challenges in addressing AI's environmental impact, including the need for cost-effective energy-saving technologies, metrics to evaluate environmental effects beyond energy consumption (e.g. water use and material impact), and the integration of technical standards into practical applications.
- b. Innovations for the energy efficiency of AI
 - Scientists are exploring ways to reduce AI's environmental impact by improving software efficiency, such as compressing large foundation models, repurposing them for multiple tasks, and employing techniques like quantization, while also addressing energy-intensive digital hardware through innovations like neuromorphic computing.
 - To reduce AI's environmental impact, the industry can focus on three key areas: adopting green energy supply, implementing energy-saving technologies and efficient products during AI operations, and promoting circularity in energy and materials to minimize emissions throughout the AI system lifecycle.
- c. Areas where collaboration and standards are needed:
 - A collaborative effort to rethink the integration of hardware and software is essential for creating sustainable AI systems, with energy efficiency being prioritized as a core value to drive environmentally friendly advancements in both scientific and industrial AI development.
 - A global, interdisciplinary effort is essential to maximize AI's environmental benefits, requiring human-centered metrics, reliable methodologies, and risk assessment to ensure effective and sustainable solutions.
 - Accurate GHG emission management, supported by relevant standards and technology solutions, is essential to identify and address critical areas for emission reduction, ensuring a systematic approach to decreasing AI's environmental footprint. Existing standards provide a foundation for sustainable AI systems, including green data centres, requirements for telecom sites, monitoring network intensity, and circular economy practices. However, as technologies advance, there is a need to update requirements continuously and further improve the design and efficiency of ICT infrastructure.
- d. The following reports on energy efficiency and AI were announced during the session:

- i. [Measuring what matters: How to assess AI's environmental impact](#) published by ITU
- ii. [Methodology to assess Net Zero progress in cities](#) – published by United for Smart Sustainable Cities (U4SSC) initiative
- iii. [Guidelines for cities to achieve carbon Net Zero through digital transformation](#) – published by U4SSC initiative
- iv. [Artificial Intelligence for Climate Action: Advancing Mitigation and Adaptation in Developing Countries](#) – published by United Nations Climate Change – Technology Executive Committee (UNFCCC-TEC)
- v. [Smarter, Smaller, Stronger: Resource-Efficient Generative AI & the Future of Digital Transformation](#) – published by United Nations Educational, Scientific and Cultural Organization (UNESCO)



Figure 29: Left: Qi Shuguang, Vice Deputy Engineer, China Telecommunication Technology Labs (systems), China Academy of Information and Communications Technology (CAICT). Right: Gitta Kutyniok, Bavarian AI Chair for Mathematical Foundations of Artificial Intelligence, Ludwig-Maximilians Universität München

5.11 Elevating AI skills for all

Naria Santa Lucia, General Manager for Skills for Social Impact at Microsoft, presented key insights into the changing workforce and the critical role that AI skilling plays in preparing workers and leaders for the new AI economy.

The main points highlighted were:

- The questions below guide Microsoft's philanthropic investments, resource allocation, and workforce development efforts in the AI economy.
 - Are we building machines smarter than people or machines that help people become smarter?
 - Are we creating machines to outperform people or to help people pursue better jobs?
- Rapid technological advancements: By 2030, 70% of job skills will change, yet 75% of young people in lower-income countries lack the skills for future jobs.
- Education gaps: Only 10% of schools and universities offer AI-related guidance despite 76% of global education leaders identifying AI literacy as essential.
- Microsoft's AI Skills Academy focuses on providing in-demand AI credentials for learners. It offers free training, curriculum, and tools like the AI Skills Navigator to help learners chart personalized learning paths. Emphasis is placed on testing and credentialing to provide learners with tangible, recognizable qualifications like digital certificates or micro-credentials.
- Microsoft is committed to democratizing AI skills globally through philanthropy, partnerships, and skilling initiatives.
- Microsoft announced a new initiative, Microsoft Elevate, to unite philanthropic efforts, sales, technical support, education, and nonprofit collaborations. The target is to train 20 million people in AI over the next two years.
- The importance of partnerships, whether through sharing resources, networks, or simply a willingness to contribute to AI skilling efforts, was discussed. This initiative is part of a broader commitment to help people shape the future of work, not just react to it. Strategic partnerships for AI skilling between Microsoft and the following organizations were mentioned for this new Initiative:
 - AI Skills Coalition with ITU and partners:
 - Encourages multi-sector collaboration to train people on AI skills.
 - Over 50 organizations have joined the coalition and it continues welcoming support.
 - UNICEF:
 - Passport to Learning: Online curriculum for low-connectivity areas, initially developed pre-COVID and used widely during the pandemic.
 - Passport to Earning: Aims to train youth in digital and AI skills, with 14 million individuals trained across 47 countries.
 - International Labour Organization (ILO):

- Women in Digital Business Program: Empowers women micro-entrepreneurs with digital and AI skills.
- Over the next five years, Microsoft Elevate will donate \$4 billion in cash and technology to schools and nonprofits.



Figure 30: Naria Santa Lucia, General Manager, Digital Inclusion and Community Engagement, Microsoft Philanthropies, Microsoft

Part 2: Thematic AI Standards Workshops

6 Trustworthy AI testing and validation

The main objectives of the Trustworthy AI testing and validation workshop were:

- a) Discuss the research around AI system testing and verification methods
- b) Provide an overview of the different methodologies that are used to test and verify AI systems and their strengths/limitations
- c) Identify any gaps in current methodologies for AI system testing and verification
- d) Explore examples of some of the methodologies and their applications in AI system testing such as Agentic AI testing and LLM security testing
- e) Discuss opportunities for international collaboration on AI testing and verification through an international collaborative platform

Collaboration will be key in developing a shared understanding of what constitutes trustworthy AI and sharing lessons learnt when it comes to best practices and appropriate technical tools and standards for AI validation and verification. The workshop's main aims were to provide information about the research trends on AI system testing and verification, covering key methods, their strengths, and limitations and the opportunities for international collaboration on AI testing.

6.1 AI system testing

The first session discussed the challenges of AI testing and current research work underway in the field of trustworthy AI testing.

Princeton University shared their work on testing autonomous driving. Trust can be placed in AI just as it is in humans. AI becomes trustworthy when models deliver consistent, error-free responses across different environments and make reliable decisions. When users see an AI system behaving predictably and dependably, they begin to trust it – just as they would a reliable person. The first and most critical step towards building trust is ensuring that an AI performs reliably even when faced with unfamiliar data. It is important that AI not only functions in controlled lab settings but also delivers consistent results when applied to real-world data. All too often, we see AI models failing to meet expectations when exposed to real-life conditions, and that undermines trust.

To assess the trustworthiness of an AI system, it needs to be examined from the perspectives of different stakeholders and its context or socio-technical ecosystem. This socio-technical "systems view" can help to understand the expected behaviour of the AI system for various input scenarios.

For example, in the case of autonomous driving, how should adequate test metrics be defined for AI system testing and what are the various contexts to be taken into consideration for AI safety?. This is a difficult and multi-faceted question, requiring conscious intervention at every step in the process of creating an AI system, from data-set collection through to deployment (and beyond). A key aspect of this process involves developing methodologies for constructing

and training AI systems that are safe or ensuring their alignment with specific contexts and scenarios.

Berkeley University presented the challenges posed by AI security risks and the research on trustworthiness and risk assessment of different Large Language Models (LLMs), which identified different ways today's LLMs can be attacked. It was noted that, with the enhanced capabilities brought about by Agentic AI systems, there is a need for collaboration among scientific communities to share lessons learned and best practices.

The Netherlands Organisation for Applied Scientific Research (TNO) has been conducting research on AI security, working closely with various Dutch government bodies and NATO partners. As AI becomes increasingly integrated into critical systems, ensuring its security is not just a technical challenge but imperative for national safety. Despite the growing importance of this field, they observed a significant gap: the lack of openly accessible tools to assess and enhance the security of AI systems. TNO is working on the development of an AI Security Assessment Framework tailored to the needs of the Dutch government. This framework aims to provide a structured approach to evaluating the security and trustworthiness of AI systems. Their findings underscore a clear call to action: the AI community must collaborate to create and share more open, practical tools that support the secure deployment of AI.

Atlas Computing presented the research concept of Flexible Hardware Enabled Guarantees (FlexHEG). As AI advances, so does the potential for catastrophic risks resulting from accidents, misuse or loss of control over dangerous capabilities. For example, severe misuse in domains such as disinformation or cyber attacks seems plausible within the next few years. As such, governance of AI technology — whether by national governments, industry self-governance, intergovernmental agreements, or all three — is a crucial capacity for humanity to develop, and quickly.

Hardware-enabled governance has emerged as a promising pathway to help mitigate such risks and increase product trustworthiness by implementing safety measures directly onto high-performance computing chips. These hardware-enabled mechanisms could be added to powerful AI accelerators used in datacentres.

However, it is not yet clear which compliance rules will be most appropriate in the future. Therefore, these hardware-enabled governance mechanisms could be considered for the flexible updating of compliance rules through a multilateral, cryptographically secure input channel, without needing to retool the hardware. Concepts like FlexHEG could enable multilateral control over AI technology, thus making it possible for a range of stakeholders to agree on a variety of potential rules, from safety rules to robust benefit-sharing agreements. Mutually agreed-upon rules could be set and updated through a multilateral and cryptographically secure mechanism to guarantee that only agreed-upon rules are applied.

6.2 International collaboration on AI testing

The session discussed the gaps in testing AI systems, evaluation of trust in AI systems, and opportunities for international collaboration to ensure the effective design, development, and deployment of AI systems that integrate considerations for AI trust. There was general agreement among participants on the need for better collaboration internationally on methodologies and scenarios for AI testing. The outcomes of this session served as the basis for the Open Dialogue on Trustworthy AI Testing taking place on 9 July.

AI technologies present opportunities that include revolutionizing science as well as risks arising from democratizing the ability to do harm, potential loss of control, or unreliable systems in domains like healthcare. Well-coordinated standards development will be key in realizing AI's benefits while guarding against risks and there is also a need to find a faster pipeline from research to pre-standardization and standardization.

As part of the G7 Hiroshima AI Process², the G7 had launched a voluntary Reporting Framework³ to encourage transparency and accountability among organizations developing advanced AI systems. The main objectives of the Hiroshima AI Process were to promote:

- A standardized way to track implementation
- Transparency and comparability
- Alignment with international AI governance initiatives

The framework aims to facilitate transparency and comparability of risk-mitigation measures and contribute to identifying and disseminating good practices. The OECD supported the G7 in developing this reporting framework to facilitate the application of the Hiroshima AI Process International Code of Conduct for organizations developing advanced AI systems⁴.

Clear quality metrics for AI testing is very important to be able to compare and assess AI systems and models. As an example, consider the current situation regarding AI testing as being analogous to comparing car models by collecting enough information to make a choice. Imagine your frustration if every conversation with a car salesperson goes like this:

You: *How fast is this car?*

Dealer: *It is really a great car. In fact, it is probably the fastest. All our customers are happy with how fast it is.*

You: *Can you give me a few numbers?*

Dealer: *Trust me, the car gets you to your destination in no time at all. It is that awesome!*

You: *Hmm - I also care about safety. What does the car offer in terms of safety?*

² See <https://digital-strategy.ec.europa.eu/en/library/g7-leaders-statement-hiroshima-ai-process>

³ See <https://transparency.oecd.ai/>

⁴ See https://www.soumu.go.jp/hiroshimaaiprocess/pdf/document05_en.pdf

Dealer: *It is a fantastic car. It is very very safe. It fulfils all the regulations. I can show you the certifications to prove it meets all the regulations. So it is completely safe.*

You: *What about the cost of running the car? How much gas does it need? What are the insurance rates?*

Dealer: *It is really cheap to run. You can trust me that it is very, very cheap. I can't give you any figures but you will be amazed how cheap it is.*

For buying a car, this is certainly farfetched. We fully expect to make our choice based on clear, understandable indicators of quality such as fuel consumption, top speed, acceleration, noise level, space for passengers, space for luggage, safety rating, stopping distance, theft protection, and charging speed. Note that this is a mix of performance, safety, security and convenience indicators.

We also fully expect that these indicators of quality can be, and have been, independently measured – be it by regulators, or by product testing organisations. We also fully expect that these indicators of quality do not just reiterate compliance with regulations (after all, we assume there will not be an illegal car on offer at the dealership), but rather measure characteristics that go beyond the scope of what is addressed by regulation.

Contrast this with procuring an AI solution like a customer service chatbot handling insurance claims or an automated landing software for a professional drone. In these cases, metrics are much harder to come by. Where they exist, they usually refer to low-level technical capabilities of the actual AI model rather than the characteristics of the whole system or solution. In many cases, they do not exist at all. Procurement departments are faced with weighing little more than marketing prose when comparing multiple solutions. Even before, they are struggling to write clear quality measures into requests for quotation because there are no widely agreed ways of measuring quality for AI solutions.

There is a closely related challenge: Suppose you are in the lucky position of having an in-house AI engineering team. When should this team stop testing and refining an AI system? When is the system “good enough” to be deployed? In fact, what does “good” even mean, and how can you measure it?

To construct a useful set of quality indicators for when an AI system or solution is good enough or which of two AI solutions is preferable, consider that a set of quality indicators might be called a “quality index”. To develop this quality index for specific use cases – comprising different indicators to evaluate different context and performance aspects, and compatible with common AI governance frameworks – a common shared language is required to be able to compare different AI solutions, or to decide when AI is “good enough” to go live.

AI systems increasingly operate in high-stakes domains, there is a need for rigorous testing methodologies that extend beyond conventional software validation to address ethical, fairness, and safety concerns intrinsic to machine learning models. The rapid evolution of generative AI and autonomous agent systems introduces complex ethical questions related to challenges including model unpredictability, bias amplification, and opaque decision pathways. Key testing techniques include adversarial testing for generative AI models that uncover

vulnerability to prompt injection and hallucination, bias auditing frameworks in agent behaviours, and explainability methods. The integration of Human-in-the-Loop (HITL) frameworks was emphasized as a dynamic control layer—examples include human-validated reward models in reinforcement learning agents and interactive prompt refinement cycles to keep generative AI outputs within ethical boundaries.

Industry could implement AI governance frameworks incorporating continuous monitoring for model versioning and explainability, coupled with automated compliance checks based on regulatory standards such as GDPR and the EU AI Act. A key focus is the integration of HITL processes, which enable real-time human oversight in AI decision-making workflows to correct erroneous outputs, enforce domain-specific ethical constraints, and adapt to evolving operational contexts.

Trustworthy AI testing needs to move beyond gatekeeper-style process checks focused solely on model performance. Instead, we could validate core assumptions, conceptual soundness, and context-aware behaviour (contextual performance), taking into account societal factors and the specific communities affected.

Establishing standards can help bridge the gap between technical metrics and societal values. Currently, there is no clear international consensus on the definitions and distinctions among AI agents, agentic AI, physical AI, and embodied AI, particularly from the perspective of standardization.

AI testing and verification are evolving beyond model capability performance assessment. Broader, dynamic scopes are now needed from both technical and socio-technical perspectives.

The sharing of best practices in AI testing is very important at the international level. Experience-based practices alone are insufficient to deal with unknown risks. A call to share logic-based frameworks across the AI ecosystem was made during the workshop. The context of AI risk assessment has shifted significantly. Key questions going forward:

- What specific technical and socio-technical AI safety risks must be addressed?
- How should these risks be addressed effectively?

The scope of AI risk management is becoming increasingly complex due to new technologies such as:

- Agentic AI
- Physical and embodied AI
- Multimodal foundation models

Adopting a "risk-chain model" looks to be essential to understand and address the complex interactions between AI systems. There is a need for international collaboration on AI testing to share best practices and assess areas where standards are needed.

User trust in AI systems is vital for their acceptance. Trust comprises elements such as interpretability, fairness, and reliability, with users' perceptions affecting their trust levels. Evaluating trust is challenging due to the need to quantify these subjective elements

consistently across various users and contexts. Current AI testing methods face limitations, including ineffective evaluation frameworks for real-world applications. Many systems excel in labs but struggle in actual deployments. The lack of standardized processes hinders comparisons across organizations, undermining AI credibility. Some of the challenges include:

- a) Lack of unified standards
- b) Lack of comprehensive testing frameworks
- c) Transparency and explainability
- d) Regulatory differences
- e) Data privacy concerns
- f) Rapid technological advancements
- g) Resource limitations
- h) Bias and fairness

International collaboration is vital for creating global AI testing standards, fostering knowledge sharing and resource integration. Collaborative frameworks can help countries develop best practices and boost AI credibility and reliability while promoting innovation and sustainable development.

The main objective of international collaboration on AI testing and assurance would be to establish unified standards for the reliable, safe, and ethical use of AI technologies. This includes developing consistent testing methodologies and fostering cooperation among governments, industry, and academia to share best practices. By addressing regulatory differences and building capacity, the goal is to create a transparent AI ecosystem that enhances public trust and mitigates risks.

An example of how countries' different risk thresholds can also impact joint AI testing was shared in relation to joint AI testing held by Japan and the UK on making LLMs reliable in different linguistic environments through assessing if guardrails hold up in non-English settings. As AI agents are being deployed globally, it is also important that these agents handle different languages accurately and consider different cultures appropriately, securely and accurately. Another challenge is the reproducibility of AI tests, which is an important objective for joint AI testing.

Common definitions for terms are needed for consistency. It was noted that ISO, the US National Institute of Standards and Technology (NIST), ITU, IEEE, OECD, and many others already provide rich vocabularies for risk, control, and evidence. The challenge is that these vocabularies only partially overlap. Translating among them is slow and error-prone, making test results hard to compare across borders or sectors. To address this, one near-term goal could be to map terminology from existing standards into a single, open glossary and invite standards bodies to validate and refine the mapping. Dozens of benchmarks exist, but most focus on accuracy or speed—not on edge-case safety, long-horizon planning, or social influence, for example. Multiple labs have begun running deeper safety evaluations, yet their protocols are rarely interoperable.

Models now update weekly, meaning certificates issued annually quickly go stale. Pioneering groups are experimenting with rolling audits and red-teaming pipelines, but the data seldom flows beyond the organization that generated it.

A new Open Alignment Assurance Initiative led by the International Association for Safe and Ethical AI aims to connect and extend these efforts rather than start from scratch. This can be addressed by linking existing academic, corporate, and national labs under a common protocol, and adding missing additional capacity as required for a test run in Nairobi to carry the same weight in New York, for example.

Some of the key takeaways on the importance of international collaboration for AI Testing are:

1. Standardized quality metrics and consistent definitions are required for testing AI systems
2. Standards on AI testing can help to support policy and legislation on AI governance
3. Sharing best practices on AI testing methods, tools, and capacity building on AI testing methodologies is essential
4. Reproducibility of AI tests is an important objective for joint AI testing
5. Identify standards gaps and initiate pre-standardization work on AI testing
6. ITU could play a leading role in facilitating multistakeholder international collaboration on trustworthy AI testing. The collaboration could focus on three key areas:
 - i. Capacity building
 - ii. Promoting standards and best practices
 - iii. Institutional frameworks.

7 Open dialogue on trustworthy AI testing

The Open Dialogue on Trustworthy AI Testing workshop brought together global stakeholders to address critical gaps and collaboration opportunities in trustworthy AI testing. The workshop employed an interactive format to facilitate focused discussions across three key pillars of AI testing collaboration: capacity building, standards and best practices, and institutional frameworks.

The main objectives of the workshop were to:

1. Identify current gaps in global AI testing capabilities
2. Explore collaboration mechanisms for capacity building in AI testing
3. Discuss standards, best practices, and conformity assessment needs
4. Examine institutional frameworks for international coordination on AI testing

The audience was divided into three thematic groups:

Group 1: Capacity Building

Focus: Understanding global capacity needs and gaps for trustworthy AI testing

Group 2: Standards, Best Practices and Conformity Assessment

Focus: Current best practices, methodologies, standards gaps, and knowledge sharing mechanisms

Group 3: Institutional Frameworks

Focus: AI governance structures and coordination at the international level

7.1 Outcomes

The outcomes of the discussions of each group are summarised below.

7.1.1 Group 1: Capacity building

The capacity building group identified several critical areas requiring attention:

Current gaps in AI testing capabilities globally

- Terminology around testing remains highly confusing and varies significantly across different contexts and regions
- Substantial disparities exist between current use-case testing and real-world testing scenarios
- Limited and unequal access to AI models, which is essential for comprehensive testing

Knowledge areas and institutional capabilities needing development

- Standardized terminologies and metrics for AI testing
 - Clear definition of roles for different stakeholders in the testing ecosystem
- Institutional capacity, which varies dramatically across different regions, particularly affecting emerging economies' ability to conduct AI testing for various use cases

7.1.2 Group 2: Standards, Best Practices and Conformity Assessment

This group focused on technical and methodological aspects of AI testing collaboration:

Priorities and gaps for collaboration

- a) Technical aspects of model testing including testing environment, data requirements, and reproducibility standards
- b) Comprehensive risk assessment covering misuse risks, AI-cyber intersections, AI-bio risks, and broader socio-technical challenges

Approaches to close the gaps

- a) Development of technical reports on standards mapping for trustworthy AI testing

- b) Establishment of pre-standardization discussions, potentially leveraging ITU platforms (e.g ITU-T Focus Groups).

7.1.3 Group 3: Institutional frameworks

The institutional frameworks group examined governance and coordination mechanisms:

Key institutional gaps

- a) Need for identifying areas requiring global alignment in AI testing approaches
- b) Requirements for agile governance structures that can adapt quickly to technological developments
- c) Strategic information-sharing, knowledge-building, and co-production mechanisms are currently inadequate

Solutions and coordination mechanisms

- a) Establishing comprehensive coordination frameworks with effective feedback loops
- b) Emphasizing complementary roles among international organizations like ITU while avoiding duplication of efforts

7.2 Future directions

Where do we go from here? Participants agreed to continue the dialogue on collaboration for trustworthy AI testing initiated at the AI for Good Global Summit to discuss how to enact some of the proposals made. Some of the key actions proposed were:

- a) **Continued dialogue on trustworthy AI testing:** Establishing a regular dialogue among the various stakeholders involved in the event was emphasized as an important need in the space, with Group 2 highlighting the need for dialogue focused on frontier model security testing and Group 1 emphasizing the need for collaboration on capacity building to enable AI testing across different jurisdictions.
- b) **Develop technical reports:** Working towards technical reports on topics related to trustworthy AI testing, such as testing environments, protocols, and risk management frameworks, was highlighted as a valuable next step. Group 3 raised the importance of strategic information-sharing, knowledge-building, and co-production mechanisms for greater institutional capacity around the world.

8 Enabling AI for health innovation and access

This section is based on the outcomes of the workshop organized by ITU, the World Health Organization (WHO), and the World Intellectual Property Organization (WIPO-), the founding UN agencies of [the Global Initiative on AI for Health](#) (GI-AI4H) during the International AI Standards Exchange. The aims of the workshop were to promote AI4H-related standardized guidelines, catalyse cross-sector collaboration, and encourage broader participation from the global health and AI communities. Designed for policymakers, technologists, health

practitioners, and humanitarian leaders, the session consisted of three themes, namely the global landscape on AI for health, frontline use cases of AI in health, and the interplay of intellectual property and AI for health.

The workshop highlighted the following:

- **Strategy and frameworks for AI in health**, highlighting the role of standards, policies, and governance in responsible AI integration in health systems, elaborating how GI-AI4H drives global cooperation with standardized frameworks, knowledge sharing, and pilot initiatives.
- **Real-world AI-enabled solutions** for global health challenges, such as LLMs for conflict-zone triage, AI-driven diagnostics for non-communicable diseases, intellectual property (IP) commercialization pathways for health technologies, etc.
- **Intersection of AI and IP law**, discussing the evolving legal landscape, ethical considerations, and strategies for managing IP in the age of AI.

8.1 Global Initiative on AI for Health (GI-AI4H)

In recognition of the importance of digital health, the rapid development of AI, the need for guidance to all stakeholders to increase adoption of AI solutions, and the pioneering work done by the ITU/WHO Focus Group on AI for Health, the leadership of ITU, the World Health Organization (WHO), and the World Intellectual Property Organization (WIPO) announced during the AI for Good Global Summit 2023 the launch of the [Global Initiative on AI for Health \(GI-AI4H\)](#).

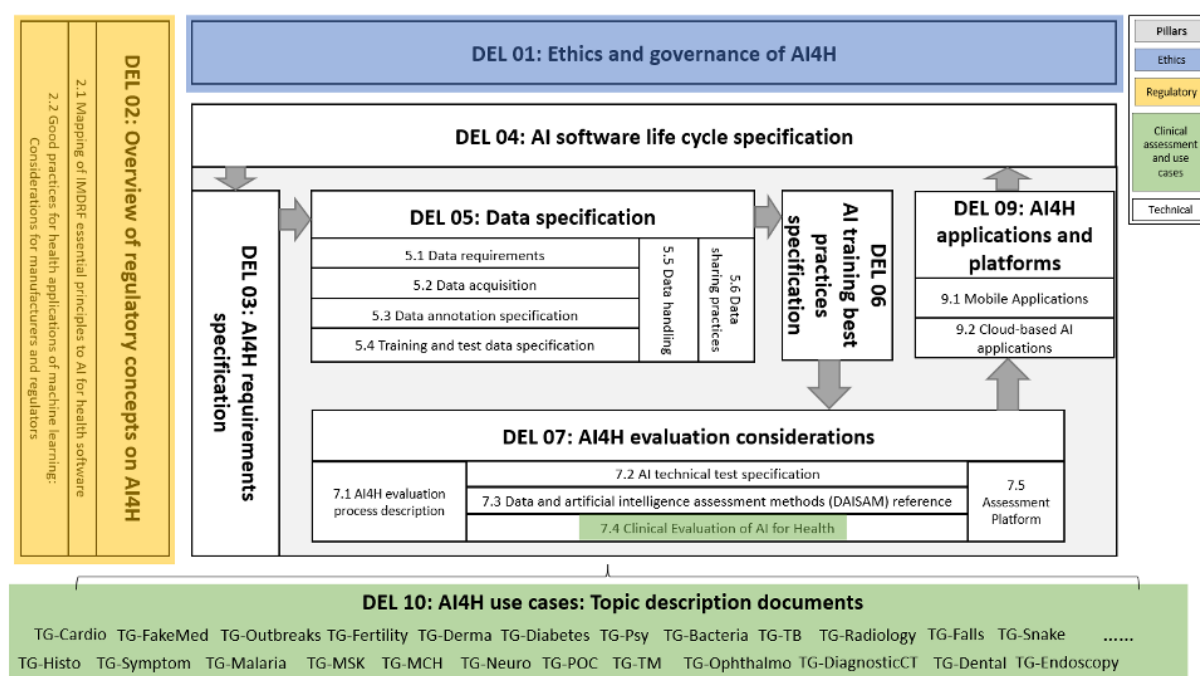


Figure 31: FG-AI4H benchmarking framework with 36 deliverables

GI-AI4H is a collaborative project to harness the power of AI for global health improvement. The initiative aims to develop international standards, norms, and policies to guide the ethical and responsible use of AI in health, promote knowledge and data sharing, foster collaboration, and build a global community of AI for Health experts. Building upon the framework with 36 deliverables from the FG-AI4H, the Initiative also seeks to establish sustainable models for implementing AI programmes at the country level, making AI solutions accessible and impactful across diverse health systems, leveraging three pillars:

1. Enable:
 - Norms, guidance, standards (including all benchmarking framework deliverables)
 - Governance (ethics, regulations)
 - Surveillance, research and evidence
2. Facilitate:
 - Knowledge sharing between countries
 - Pool funding
 - Cooperation between all stakeholders
3. Implement:
 - Scale programme in countries (phase 1 to target 12 to 18 countries)
 - Build capacity building for AI for health programmes
 - Build sustainability models

8.2 Strategy and Frameworks for AI in health

On the policy and governance challenges and opportunities associated with the adoption of strategic frameworks for AI in health, topics discussed ranged from regulatory harmonization, ethical implementation, and equitable access to how GI-AI4H can facilitate global collaboration.

The following challenges in implementing AI in health were highlighted by speakers during this session:

- a) AI for health must be customized to local needs and infrastructure. Meanwhile, it should evolve to ensure transparency, accountability, and ethical oversight.
- b) The importance of localization and sovereignty of LLMs based on experience in developing contextualized, trustworthy, and transparent LLMs.
- c) The EU's efforts to tackle medical data management and its proposed principle-based approach enabling agile adaptation to technological changes.
- d) There is a need to build communities and champion networks to guide AI deployment.
- e) AI should be viewed as an enabler, not a replacement, in health.
- f) Global collaboration is vital to share best practices and avoid duplication.
- g) Investments are needed in research, infrastructure, regulation, and human capacity.
- h) Human-centred validation of AI models is essential for trust and adoption.
- i) Standards are needed for the validation of models and must support transparency, reproducibility, and contextual validation.

The following action areas were identified for strategy and frameworks for AI in health:

1. Develop agile, standards-based regulatory templates to help countries adapt AI frameworks to local health needs while ensuring global alignment.
2. Expand the global AI Champions network to foster cross-border collaboration and coordinate research, policy, and implementation efforts in AI for health.
3. Support platforms enabling benchmarking and localized AI validation not requiring AI expertise, helping ensure technologies are contextually relevant, trustworthy, and aligned with regional priorities.
4. Create model laws and standardized frameworks to guide national AI health strategies, promoting transparency, ethics, and regulatory coherence across jurisdictions.
5. Advocate for open science and foster inclusive, transparent, and accountable AI policy and research that benefits society as a whole.

8.3 Use cases of AI in health

Some real-world applications of AI to address global health challenges were showcased at the workshop, highlighting the transformative impact of AI in the health space and some of the explored the necessary factors for key stakeholders to bring AI to the people at the frontline.

The following applications of AI in health were presented:

- *WHO and Training Aid* presented how AI facilitates the development and customization of standards-guided training packages for emergency agencies. These packages scale up the learning and competencies of emergency workers in different roles at the country level.
- *Peking University, China* presented China's "5-minute social rescue circle" initiative for cardiac arrest emergencies to improve survival rates, demonstrating how AI, data integration, and community-based approaches could increase bystander CPR and survival rates.

Box 1: AI impact on healthcare

The advancements in AI technology present unprecedented opportunities to revolutionize healthcare, making it more effective, accessible, and economically sustainable. By fostering the integration of AI through appropriate policies, countries can improve healthcare and ensure that new technologies, treatments, and medicines benefit society at large. For example:

- AI can facilitate the efficient allocation of healthcare resources. Predictive modelling can forecast patient admissions and optimize the use of hospital beds, staff, and equipment. This helps ensure that resources are available where and when needed most, reducing waste and enhancing the quality of care.
- AI also has the potential to tackle some of the most pressing challenges in healthcare, such as rising costs, inefficiencies, and the demand for higher-quality care.
- AI can reduce costs and streamline administrative tasks like patient scheduling, billing, and electronic health records management by automating and optimizing operations. This can free healthcare professionals to focus more on patient care.
- In diagnostics, AI enhances accuracy and enables earlier detection, often leading to less invasive and more cost-effective treatment options.
- AI-driven personalized treatment plans can complement traditional approaches by offering more targeted and effective care, improving patient outcomes while also helping to reduce the financial burden on healthcare systems.

- *RevolutionAIze* shared a public health tool from India that uses AI to track and address childhood malnutrition. The tool analyses existing WHO standards and integrates with government healthcare applications.
- *Uppsala University, Sweden and University of Helsinki, Finland* presented an AI-based point-of-care diagnostics, which have also been applied in East Africa for cancer screening and infectious disease diagnostics using mobile microscopy and cloud-based AI systems.
- *WHO Eastern Mediterranean Regional Office* shared the AI-powered All Hazard Information Management toolkit, which speeds up outbreak decision-making and reduces document preparation time.

The following key takeaways were seen as important factors for success:

- AI tools must be localized, culturally appropriate, and usable offline to ensure equity.
- Partnerships with governments, non-governmental organizations (NGOs), and academia are critical for scale and sustainability.
- AI solutions should be designed for low-resource settings and integrated into existing systems.

The following follow-up action areas were identified:

- a) Develop standards for AI validation and contextual adaptation.
- b) Support open, modular platforms for AI deployment in public health.
- c) Advocate for equitable access and sustainability of AI tools.
- d) Evolve a global experts network and develop partnerships for responsible use of AI.
- e) Encourage data governance frameworks that support local ownership and privacy.

8.4 IP management to enable AI implementation for healthcare innovation and access

This session, hosted by WIPO's External Relations and Global Challenges Divisions, explored how private innovators and governments can effectively leverage IP to ensure AI health solutions reach markets and expand healthcare accessibility in low- and middle-income (LMICs) countries.

The following main points were highlighted when it comes to IP management:

- Africa needs to establish autonomous AI legal frameworks for healthcare, prioritizing people-centred, ethical regulation amid fragmented governance. IP can incentivize innovation but must be balanced with access and equity, especially in LMICs.
- Copyright management was highlighted as central to deploying AI in health, emphasizing that fair use, transformative use, and derivative works are key legal concepts in AI.
- The industry uses AI in drug discovery and clinical trials but often keeps tools proprietary due to data privacy. IP frameworks must align with AI's technical realities for cross-boundary functionality.

Box 2: Importance of IP for AI in healthcare

Medical technology is revolutionizing health care. From AI-driven diagnostics to 3D-printed implants, innovative solutions are enhancing patient care but also highlighting the critical role of intellectual property (IP) in bringing such advances to market. With AI and data analytics now integral to health care innovation, the importance of protecting IP in this domain has never been greater. For companies, data is not just a by-product of their technology, it is a vital asset that drives the effectiveness of AI models and opens up new avenues for licensing and collaboration.

Brain scanners, for example, relies on AI models trained with diverse datasets to accurately detect brain injuries. Protecting that data through robust IP strategies is crucial not only to protecting the innovation but also to capitalizing on its full market potential. For medical technology companies, securing IP is only part of the challenge. Bringing a medical device to market involves navigating complex regulatory landscapes, which vary significantly across different jurisdictions. IP plays an important role in that process, providing a foundation for compliance and market access.

In terms of key areas of action, the following were proposed:

- a) IP frameworks must evolve to address AI-specific challenges like authorship and data rights.
- b) Exceptions and limitations are essential to enable AI development.
- c) Open innovation and indigenous data governance must be considered in global standards.
- d) Global cooperation is needed to align IP, data, and regulatory frameworks.
- e) Develop model IP clauses and licensing templates for AI in health.
- f) Promote awareness of copyright implications at each stage of the AI pipeline.
- g) Encourage countries to include IP in their AI strategies and policies.

8.5 Technical brief on AI in traditional medicine



WHO's Global Traditional Medicine Center presented its efforts to bring evidence-based traditional medicine to billions worldwide and the related ITU/WHO/WIPO activities under the GI-AI4H. The challenges and opportunities of integrating AI with traditional medicine are described in the technical brief, which also covers the development of a global library on traditional medicine featuring an AI interface. Presenting at the summit, WHO highlighted the importance of respecting traditional knowledge and biodiversity while promoting innovation and access to healthcare.

The technical brief was developed by leveraging the findings of a literature review and supplementing this with knowledge and inputs captured during the conceptualization process with experts from the [Topic Group on AI and Traditional Medicine \(TG-TM\)](#) under the ITU-WHO Focus Group on Artificial

Intelligence for Health (FG-AI4H).

8.6 Closing summary

During the closing, ITU, WHO and WIPO high-level representatives emphasized the importance of ethical AI use in health, global collaboration, and equitable AI access, outlining a strategic vision for responsible AI for health. They called for collective action to develop standards, advocate sustainable adoption, foster partnerships, share best practices, invest in open science and research, and advocate inclusive, transparent, and accountable AI policy and regulation that embraces innovation and protects the vulnerable. The vision centres on an ethical, inclusive, and globally coordinated AI healthcare future that leaves no country is left behind, via activities such as those driven by the Global Initiative on AI for Health aimed at balancing innovation with accessibility, transparency, and local ownership.

9 Human-centred AI for disaster management: empowering communities through standards

9.1 AI for disaster management

AI can support early warning systems, integrate with nature-based solutions, and transform risk data into actionable insights that protect communities and save lives. The outputs of the [ITU/WMO/UNEP Focus Group on AI for Natural Disaster Management \(FG-AI4NDM\)](#) and the activities of the [Global Initiative on Resilience to Natural Hazards through AI Solutions](#) play an important role in facilitating the adoption of AI and supporting standards for disaster resilience.

The Focus Group was supported by ITU, the World Meteorological Organization (WMO), and the UN Environment Programme (UNEP), partners now joined by the UN Framework Convention on Climate Change (UNFCCC) and the Universal Postal Union (UPU) in support of the Global Initiative.

AI is revolutionizing how we predict, prepare for, and respond to disasters by enabling monitoring systems, real-time impact analysis, and optimized relief efforts. The efficacy of AI-oriented disaster management systems can be enhanced using standards that help ensure their reliability in predicting hazards. Human-centered AI solutions that address community needs, multilingual support, and digital divides are essential to build resilience in vulnerable populations.

The integration of AI in disaster management offers several advantages that enhance the overall response capability:

- **Disaster prediction:** AI-based algorithms can detect subtle signs of impending crises that might be missed by human analysts, thus providing decision-makers with vital advance notice.
- **Weather forecasting:** Sophisticated AI models can greatly improve the accuracy of weather predictions, helping to anticipate and prepare for events like hurricanes or heat waves.
- **Disaster response and recovery efforts:** AI can help streamline the coordination of response teams and the distribution of aid by quickly identifying the most critical needs and ways to address them.
- **Post-disaster recovery and rebuilding:** By analysing data on damage and resource availability, AI can assist in developing recovery plans that are both efficient and equitable.
- **Providing basic medical and psychological consultations:** Robotic process automation and AI-supported chatbots can provide immediate, around-the-clock support for basic healthcare and psychological needs in the aftermath of a crisis.

The role of AI in transforming disaster management includes enhancing preparedness, response, and resilience. Some examples include satellite-based anomaly detection, real-time crisis mapping to AI-powered public warning systems, climate forecasting, and data-driven

postal resilience strategies. At the same time, interoperability, international standards, and collaborative frameworks play an important part in ensuring that AI tools are reliable and capable of serving vulnerable communities effectively. All presentations are available [here](#).

9.2 Global collaboration on standards for disaster management

The [Global Initiative on Resilience to Natural Hazards through AI Solutions](#) is exploring AI use cases, providing expert guidance, and supporting research, innovation, and standards development amid increasing climate volatility and disaster risks worldwide such as seismic, hydrometeorological, and other natural hazards,.

It also aims to create an AI readiness framework to assess and improve national capacities for using AI in disaster management. Technical standards are key to ensuring AI is used safely, responsibly and equitably in disaster management – a field where decisions must be made quickly and carefully.

9.3 Use cases of AI for disaster management

Some examples of use cases being considered at the level of the Global Initiative on Resilience to Natural Hazards through AI Solutions include:

- High-Performance Computing enhances crisis response by accelerating big-data insights, detecting anomalies in satellite images, and generating regional crisis maps. Tools such as the [AI Digital Assistant ECHO](#) are used to build holistic, AI-driven views to support decision-making. AI4Space is used for automating Earth Observation (EO) data processing, compressing hyperspectral data, and accelerating forecast simulations to aid responders.
- A Multi-Domain Marketplace within the Europe Space Agency (ESA) Civil Security from Space Hub, integrating SatCom, EO, Internet of Things (IoT), and AI to enable faster crisis response. It showcased platforms like Unified Communication Services Manager (UCSM) for SatCom booking, CGI Sense360 for EO services, and Public Sector Information (PSI) standards to enhance interoperability. Collaboration is ongoing to build a next-generation Smart Digital Marketplace for crisis management.
- The need for unified public warning systems and improved disaster communication using geolocation, multimedia, and AI chatbots. Examples like I-REACT and SAFERS highlight the role of human-centered AI in crisis management. It underscored the need for adopting the Common Alerting Protocol standard ([CAP](#)), multi-channel communication, and user-focused AI to build resilient communities.
- UPU has developed strategies for risk management and resilience in the postal sector across its 192 member countries. It covered pre- and post-disaster measures, including the Disaster Resilience Fund, Business Continuity Plans, and the Emergency and Solidarity Fund. Key initiatives like the [Unified Data Platform \(UDP\)](#) and 2IPD were showcased to enhance data-driven decision-making and operational continuity.
- Initiatives like the [AI Climate Application Hub](#) and [AI for Climate Action Award](#) under the [UNFCCC Technology Mechanism](#) highlighted AI's role in early warning systems, disaster management, and climate forecasting, especially for developing

countries. Risks such as digital divides, data biases, and energy use were highlighted, with a call for inclusive governance, better infrastructure, and global cooperation.

- Early Warning For All (EW4A) Unconnected Demo showcased how AI uses satellite imagery and deep learning to create detailed population maps, identifying vulnerable and unconnected communities. It highlighted the importance of accurate data for humanitarian response, with examples from the Solomon Islands and flood monitoring in Ethiopia. Demonstration also introduced the Precision Populations Early Access Program and shared resources like historical flood datasets and predictive models.

9.4 How AI transforms disaster management

Real-world case studies from the ongoing activities of the [WMO-led Working Group on Digital Transformation for Hydrology and Water Resources](#) highlight the capabilities of AI-driven tools such as route optimization for responders, flood forecasting, and user-centric chatbots elevating situational awareness and decision-making are highlighted.

AI enhances disaster management and climate risk prediction through human-centred, standards-driven approaches. This is achieved by systemic views of climate impacts, spatial prediction using deep learning, and the integration of geospatial, meteorological, and socio-economic data. Key themes include early warning systems, LLMs for communication, and federated learning to build equitable, global AI solutions.

The AI tool, [SukhaRakshakAI](#) (Anticipatory Drought Intelligence for a Climate-Resilient Future) empowers farmers and drought managers with early forecasts and localized advisories. SukhaRakshakAI aims to address global drought challenges such as weak early warning systems and poor data availability, while highlighting historical impacts on agriculture. Built on prediction, preparation, and protection, the system integrates multi-source data, AI models like Gemini and AI4Bharat, and multilingual support.

The United Nations Office for the Coordination of Humanitarian Affairs (UNOCHA) works to improve data use in humanitarian response through four streams: Data Science, Responsibility, Services, and Learning. The [Humanitarian Data Exchange \(HDX\)](#) platform hosts 20,000 datasets and supports crises like COVID-19. UNOCHA leverages AI for climate forecasting, anticipatory action, and early warnings, emphasizing responsible governance and ethics.

The development of digital twin technology using satellite data, illustrated through a case study in the Kingdom of Tonga by United Nations Office for Outer Space Affairs (UNOOSA). UNOOSA plays a pivotal role in maintaining the UN Space Registry, advancing global development through space science, and providing technical expertise and training. Through its UN-SPIDER platform, UNOOSA ensures universal access to space-based information for disaster management, covering the entire disaster management cycle. The implementation of digital twin technology, particularly through high-resolution satellite imagery and advanced AI algorithms, provides a dynamic and detailed representation of the real world which helps users understand the impact of rising sea levels. This technology can also aid decision-makers in disaster preparedness and response by simulating potential disasters from rising sea levels. The Tonga Disaster Preparedness Pilot Project demonstrates how remote sensing and digital twin technology can simulate disaster scenarios to assess

potential damage and improve preparedness strategies. By combining satellite imagery and AI, digital twins create detailed, cost-effective 3D models compared to unmanned aerial vehicles (drones). UNOOSA also integrates IoT sensors with digital twin technology to optimize evacuation planning and support data-driven infrastructure planning.

For Tongatapu Island, Kingdom of Tonga, these digital twin products allow the identification of vulnerable coastal areas and contribute to planning evacuation routes, reinforcing coastal defences, and developing disaster-resilient infrastructure. The digital twin can serve as a collaborative platform for national government agencies, local communities, and disaster response teams. By enabling them to interact with the same data and models, stakeholders can plan and execute disaster response strategies more efficiently. The digital twin could also incorporate predictive analytics, helping to anticipate future risks based on current trends and historical data, making it a vital tool in long-term disaster resilience planning.

In essence, creating a digital twin with Geographic Information System software involves leveraging its full suite of tools for mapping, simulation, and analysis, while continuously integrating real-time data to reflect the current state of the environment. This helps decision-makers to use the digital twin to both prepare for and respond to disaster scenarios in a timely and effective manner.

The continuous integration of real-time data ensures the digital twin remains updated and relevant, making it a valuable tool for both immediate disaster response and long-term resilience planning.

9.5 Conclusion

1. Human-centred AI solutions that address community needs, multilingual support, and digital divides are essential to build resilience in vulnerable populations.
2. Standards are vital to interoperability, reliability in AI-driven disaster management systems, enabling data integration, and coordinated responses globally.
3. ITU plays a key role in leading global standardization efforts for disaster response and recovery, fostering meaningful partnerships including through the [Global Initiative on Resilience to Natural Hazards through AI Solutions](#) to develop guidelines and frameworks for the adoption of AI-centric disaster management solutions.
4. AI-powered monitoring systems, digital twins, and satellite data analytics can significantly enhance disaster prediction, situational awareness, and crisis response.
5. Areas where standards are needed include interoperability protocols for multi-domain data, unified communication standards for public warnings (e.g., CAP), and frameworks to ensure responsible AI use.
6. Emerging areas for standards development:
 - Digital twin, earth observation, and remote sensing technologies for reliable modelling and disaster simulations.
 - AI can play a critical role in combating slow-onset disasters like desertification and land degradation.

10 AI and multimedia authenticity standards

Seeing is believing, or is it? Today, AI-generated deepfakes, manipulated media, and algorithmic disinformation are making it harder than ever to discern fact from fiction. Misinformation and disinformation lead the short-term risks identified by the World Economic Forum Global Risks Report 2025⁵ and may fuel instability and undermine trust in governance, complicating the urgent need for cooperation to address shared crises.

The AI and multimedia authenticity standards workshop aimed work with the following objective to:

- a) Identify and analyse risks, opportunities, and challenges in multimedia authenticity in relation to AI
- b) Provide an overview of the work undertaken by the AI and Multimedia Authenticity Standards Collaboration (AMAS)
- c) Highlight the importance of policy measures for international AI governance for AI
- d) Discuss the effectiveness of AI watermarking, provenance, and deepfake detection technologies as well as their application use cases and gaps that need to be addressed
- e) Explore opportunities for collaboration on standardization activities on AI and multimedia authenticity

The workshop consisted of the following sessions:

- Overview of the AI and multimedia authenticity standards collaboration
- Fighting misinformation through fact-checking and deepfake detection
- Presenting the new AMAS paper, "Standards Mapping Landscape for AI and Multimedia Authenticity"
- Presenting the new AMAS paper, "From policy to practice: building trust in multimedia authenticity through international standards"

10.1 Overview of AI and Multimedia Authenticity Standards Collaboration

The [AI and Multimedia Authenticity Standards Collaboration](#)⁶ (AMAS) is a global initiative under the [World Standards Cooperation](#) (IEC, ISO and ITU) announced at the AI for Good Global Summit 2024.

This initiative aims to redefine digital integrity with an inclusive and future-oriented framework of transparency, accountability, and ethical innovation. By building a cohesive ecosystem of international standards, the collaboration is laying the foundations for a world where both creators and consumers can trust what they see, hear, and share.

AMAS work, furthermore, helps realize the objectives of the Global Digital Compact, adopted by the UN General Assembly as a framework for countries and industries to ensure that AI and other technologies benefit all of humanity.

⁵ See <https://www.weforum.org/publications/global-risks-report-2025/>

⁶ See <https://aiforgood.itu.int/multimedia-authenticity/>

Updates on AMAS activities and reports were provided during the workshop (see figure below).



Figure 32: Overview of AMAS work to date since set up

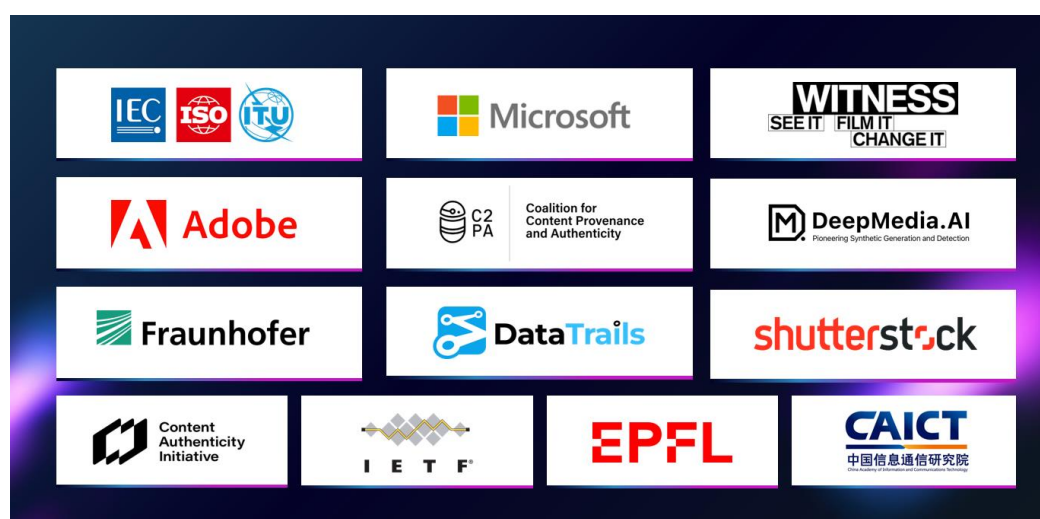


Figure 33: Organizations participating in AMAS

AMAS is focused on advancing standardization in the rapidly evolving field of AI-generated and altered media. By identifying gaps and driving the development of new standards, it supports transparent, privacy-conscious, and rights-respecting practices. AMAS also aims at informing policy and regulatory frameworks to promote legal compliance and safeguard public trust. AMAS serves as a vital forum for dialogue among standards developers, civil society organizations, technology companies, and other key players. Participating organizations (See figure above) include the International Electrotechnical Commission (IEC), the International Organization for Standardization (ISO), the International Telecommunication Union (ITU), the Coalition for Content Provenance and Authenticity (C2PA), the China Academy of Information and Communications Technology (CAICT), DataTrails, Microsoft, Deep Media, and WITNESS. The leadership team of AMAS is shown in the figure below.

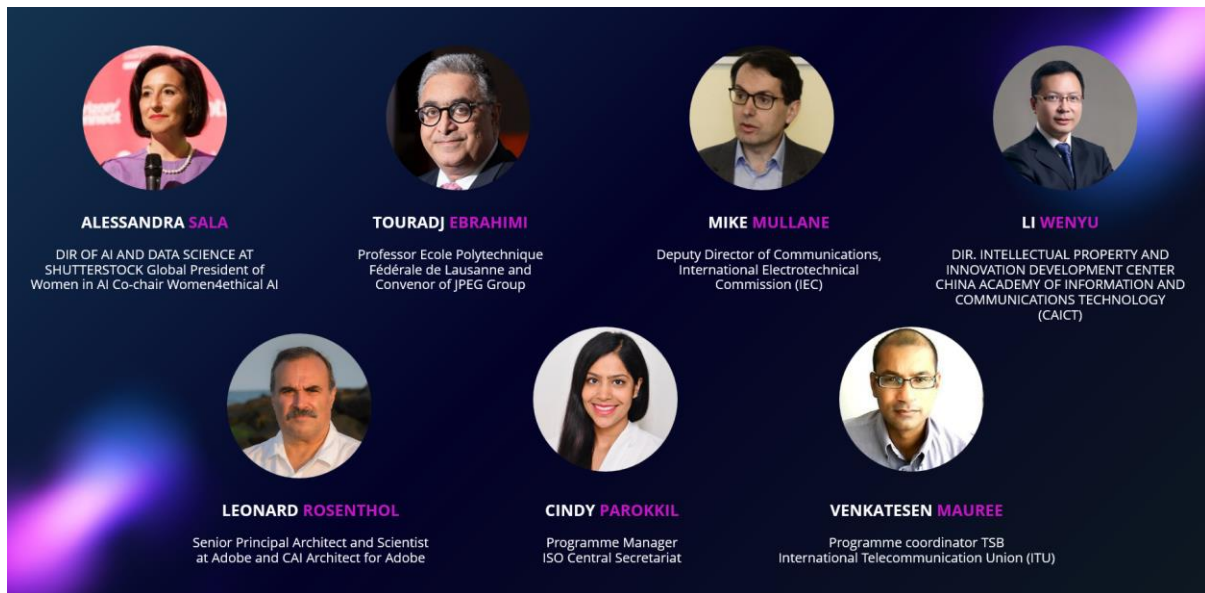


Figure 34: AI and Multimedia Authenticity Standards Collaboration leadership team

AMAS work is centred around three pillars:

- a) Technical Activities - This pillar is working on the mapping of the standardization landscape in the areas of AI and multimedia authenticity, including the identification of gaps where standards are still needed to support relevant policies across governments, industries, and other stakeholders want to implement.
- b) Policy – A forum for connecting with governments to discuss the alignment of policies and standards and related lessons learned.
- c) Communications - Promoting the work of AMAS and its value to various stakeholders (e.g, governments, social media platform providers, media companies, consumers, etc.)

Two AMAS papers were presented at the summit:

- i. Technical Paper on "[Mapping the Standardization Landscape](#)" developed by the Technical Activities pillar
- ii. Policy Paper on "[Building Trust in Multimedia Authenticity through International Standards](#)" developed by the Policy pillar

10.2 Standards mapping landscape on AI and multimedia authenticity

The technical paper on AI and Multimedia Authenticity Standards provides a comprehensive overview of the current landscape of standards and specifications related to digital media authenticity (covering different media types from images, audio, video, data, and PDF, among others). It categorizes these standards into five key clusters: content provenance, trust and authenticity, asset identifiers, rights declarations, and watermarking. The report provides a short description of each standard along with link for further details.

By mapping the contributions of various standards bodies and groups, the paper aims to identify gaps and opportunities for further standardization. Over 35 technical standards from SDOs are

analysed in the standards mapping. This could serve as a valuable resource for stakeholders seeking to navigate the complex ecosystem of standards at the intersection of AI and authenticity in digital media and to implement best practices to safeguard the authenticity of digital assets and the rights assigned to them. The findings underscore the critical role of robust specifications and standards in fostering trust and accountability in the evolving digital landscape.

As next steps, focus will be required on the harmonization of overlapping standards and the development of interoperable frameworks that can be widely adopted across industries. Emerging areas of work, such as the integration of decentralized technologies for enhanced provenance management and the exploration of new watermarking techniques for synthetic media, present exciting opportunities for innovation. Additionally, fostering awareness and adoption of these specifications and standards through education, advocacy, and pilot implementations will be crucial in ensuring their effectiveness and impact. The report is intended to be a living document and will be updated over time.

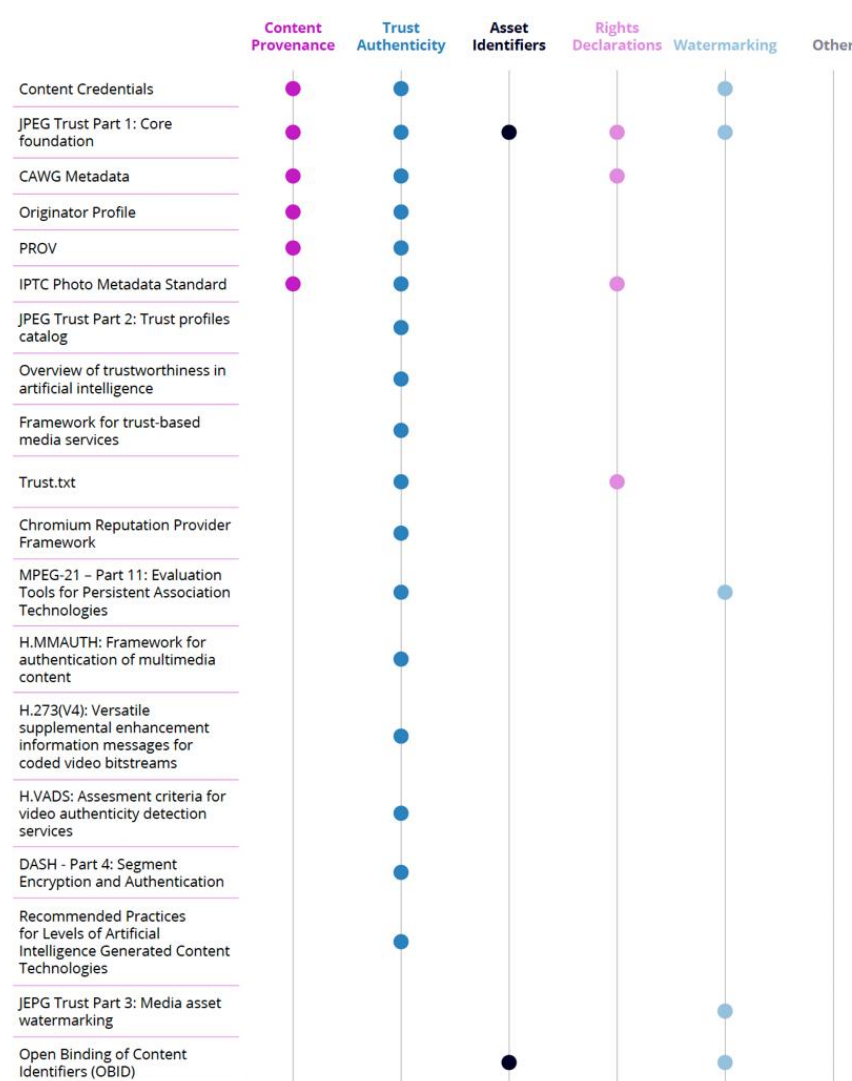


Figure 35: Snapshot of standards mapping landscape from the Technical Paper



Figure 36: Snapshot of standards mapping landscape from the Technical Paper - continued

10.3 Policy Paper: Building trust in multimedia authenticity through international standards

The Policy paper is primarily aimed at policymakers and regulators. It seeks to demystify the complexities of regulating the creation, use, and dissemination of synthetic multimedia content through prevention, detection, and response, and to present these issues in a clear and accessible manner for audiences with varying levels of expertise and technical understanding.

According to the paper, misinformation and disinformation are being used as almost interchangeable terms, but actually differ considerably, especially when it comes to their motives and application. Often overlooked in discussions is "malinformation." Malinformation, in the context of fake news for example, can be especially dangerous when

used in conjunction with disinformation as part of orchestrated campaigns intended to spread untruths. The paper provides definitions for misinformation, disinformation and malinformation as follows:

- Misinformation refers to false information but is not created or shared with the intention of causing harm.⁷
- Disinformation is false content intentionally created and disseminated to mislead, harm, or manipulate.
- Malinformation is factual information used out of context with the intent to cause harm. For example, publishing private data with malicious intent (e.g. revenge porn or non-consensual intimate imagery) or altering contextual metadata to mislead.

Types of misinformation, disinformation, and malinformation vary considerably. The table below from the paper provides some examples.

Table 1: Types of misinformation, disinformation, and malinformation (Source: Policy Paper)

Fabricated content	Usually, 100 % false and designed to deceive and do harm. ⁸ Distinguishing between the real and fabricated content is extremely difficult. Exposure to sophisticated deepfakes used to promote fabricated content can deeply impact trust in the messages citizens receive.
Manipulated content	Genuine information or imagery that has been distorted. These types of content often manipulate genuine content by doctoring an image, or use sensational headlines or click bait.
Imposter content	Impersonation of genuine sources, very often using the branding of an established agency or a reputable news agency. This form of disinformation takes advantage of the trust people have in a specific organization, a brand or even in a person. Adversaries will use phishing and smishing messages using a well-known brand in an attempt to create an impression that the recipient(s) are receiving legitimate content.
Misleading content	Misleading information is created by reframing stories in headlines. This typically uses fragments of quotes to support a wider point, often citing statistics in a way that aligns with a position. Alternatively, it can be the deliberate decision not to cover something because it undermines an argument. When making a point, everyone is prone to drawing out content that supports their overall argument.
False context	Factually accurate content combined with false contextual information, such as the headline of an article failing to reflect the content. Basically, the genuine content has been reframed. False context images are a low-tech but still a powerful form of misinformation and disinformation.
Satire and parody	Humorous but false stories passed off as true; there is no intention to harm, but readers may be fooled. What was once treated as a form of art,

⁷ See <https://webarchive.unesco.org/web/20230926213448/https://en.unesco.org/fightfakenews>, or non-consensual

⁸ This type uses false content such as the example of a deepfake audio clip of London mayor Sadiq Khan that was widely circulated on social media in November 2023. The actors used a simulation of the mayor's voice allegedly calling for pro-Palestinian marches to take precedence over Remembrance weekend commemorations on the same day.

	is now vigorously used to intentionally spread rumours and conspiracies. It is difficult to police as the perpetrators argue they are merely doing something that shouldn't be treated seriously or literally. The danger of this type of misinformation and disinformation is in the method and speed with which it gets re-shared. In doing so it is often reshaped or reframed and a wider audience loses the connection with the original messenger, failing to understand it as satire.
--	--

The report provides an overview of international guidelines, and national regulatory frameworks (such as the EU AI Act) guiding efforts to regulate online safety, misinformation, and disinformation. These involve contributions from governments, international organizations, civil society, and technology companies, often emphasizing concepts such as safety-by-design, transparency, and accountability.

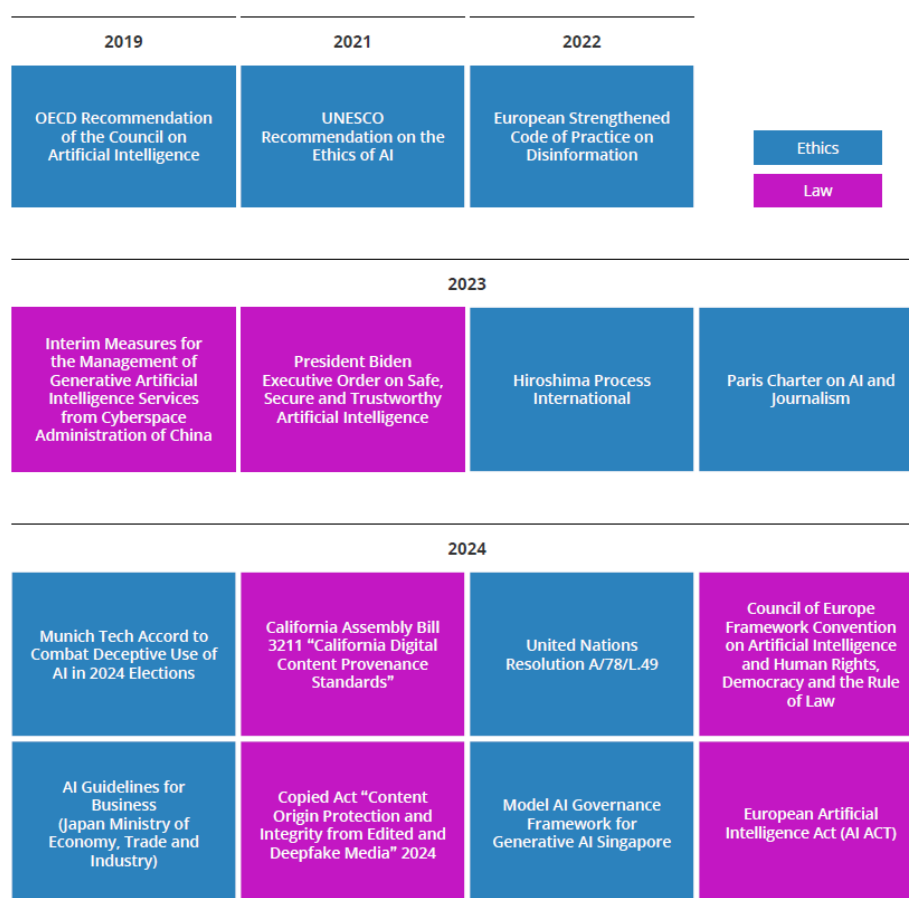


Figure 37: international guidelines, and national regulatory frameworks guiding efforts to regulate online safety, misinformation and disinformation (Source: Policy Paper of AMAS)

In addition, the paper aims to highlight global initiatives and underscore the vital role and benefits of international standards in promoting regulatory coherence, alignment, and effective enforcement across jurisdictions.

The paper offers practical guidance and actionable recommendations, including a regulatory options matrix designed to help policymakers and regulators determine what to regulate (scope), how to regulate (voluntary or mandatory mechanisms), and to what extent (level of effort). It also explores a range of supporting tools – such as standards, conformity assessment

mechanisms, and enabling technologies – that can contribute to addressing the challenges of misinformation and disinformation related to multimedia content. At the same time, it emphasizes the importance of striking a balance that enables the positive and legitimate use of either fully or partially synthetic multimedia for societal, governmental, and commercial benefit.

One of the major challenges faced by policymakers and regulators is that multimedia authenticity in the case of generative AI is fundamentally a "Black Box." There is limited transparency about how these models are developed and trained. The question that looms is how to enable effective governance when the underlying operations are largely opaque. The main challenges about how to ensure trustworthiness and interpretability of multimedia content without stifling innovation intersects with broader concerns. These include how to align with emerging global priorities, such as combatting misinformation, and how they can be shaped or influenced by online safety regulations.

According to the paper, there is a recognition from multiple stakeholders that regulatory and enforcement bodies cannot alone build trust in multimedia. All stakeholders need to work together and find new forms of international collaboration and regulation, even perhaps self-regulation. This needs to be coupled with corporate responsibility that fosters trust and promotes human rights, media literacy, and ethics.

The paper proposes the adoption of Prevent-Detect-Respond (PDR) frameworks to build trust in multimedia authenticity to address the above challenges. This three-pronged approach aims to provide a scalable, flexible structure that balances regulatory intent with technical feasibility. This framework mirrors existing approaches to privacy (e.g. GDPR and the California Consumer Privacy Act) and cybersecurity (e.g. NIST cybersecurity framework and the Payment Card Industry Data Security Standards). The strength of PDR lies in its simplicity and versatility; it is widely understood, adaptable throughout sectors, and conducive to regulatory alignment. In the case of privacy, successful approaches emphasize prevention (privacy-by-design), detection (breach notification and monitoring), and response (enforcement actions and mechanisms for user redress).

Table 2: Applying PDR framework to build trust in multimedia authenticity

Approach	Policy Requirements	Method	Benefit and/or outcome
Prevention	Transparency	Labelling	Informs users about various aspects of the content. Clearly identifying if the content was AI generated.
		Watermarking	Non-human perceptible markings applied to content that provide information about it.
	Traceability	Content provenance tools	Enables providing information about the content's origin and changes to establish accountability and attribution.
	Accountability	Conduct risk assessment	Enforcement can be made more efficient when areas are identified as high risk. Prevalent abuse or patterns of behaviour are identified and treated as priorities. This proactive approach helps mitigate the risks associated with manipulated content,

Approach	Policy Requirements	Method	Benefit and/or outcome
			ensuring that users are protected from misinformation and fraudulent activities.
	User education	Public awareness initiatives	Reduces accidental misuse through education about copyright laws and the consequences of infringement.
Detection	Detecting manipulated content and deepfakes	Technological solutions	These solutions offer numerous benefits such as protecting intellectual property, verifying image, audio, text and video authenticity, and aiding in online safety and security. However, it creates a ‘back and forth war’ with bad actors who attempt to avoid these detectors. For example: https://arxiv.org/abs/2504.2148
	Data privacy	Data handling and adherence to data protection legislations	All data processed are subject to randomized manual review, ensuring accuracy and compliance with data protection legislation.
Response	Enforcement	Regulatory interventions	Penalties can be applied and rules enforced through governments enacting laws and regulations that specifically address the techniques and approaches that should be used. They also address what happens when such techniques are breached.
	Explainability	Use of explainer-type algorithms, AI model verification methods and information about training datasets used.	Decisions made by AI systems can be checked to maintain a high level of reliability and trustworthiness. This helps mitigate risks of IPR breaches.
	Dispute mechanisms	Content contestability	Clear and well communicated mechanisms benefit individuals, helping them dispute claims.
		Platform bans	Policing of problematic areas can be more effective and beneficial when access to platforms and websites that frequently host infringing content is regularly removed.

Finally, the paper includes a set of practical checklists for use by policymakers, regulators, and technology providers. These can be used when designing regulations or enforcement frameworks, developing technological solutions, or preparing crisis response strategies. The checklists are intended to help align stakeholder expectations, identify critical gaps, support responsible innovation, and enable conformity with emerging standards and best practices.

The following recommendations are intended for policymakers, regulators, the media and technology sectors, and standards bodies. Each can be operationalized swiftly to strengthen multimedia content authenticity and build global trust.

The paper makes the following recommendations.

For policymakers and regulators:

- Consider the checklists provided in section four of the Policy Paper.
- Participate in and collaborate on standard development and alignment initiatives, especially through multilateral forums to help promote regulatory alignment.
- Consider international standards when developing and implementing regulatory sandboxes to test new technologies, policy approaches, and compliance models in controlled environments.
- Adopt a PDR framework based on internationally recognized standards to structure responses to content authenticity challenges.
- Consider data privacy and bias regulations to ensure AI-generated content respects user rights and avoids discriminatory outcomes.
- Support and encourage the development of conformity assessment frameworks specifically targeting multimedia content, incorporating requirements related to AI risks, misinformation, disinformation, and deepfakes.
- Consider a conformity assessment and/or certification scheme for multimedia content authentication based on international standards, including relevant testing.

For technology developers and providers, policymakers could request that they consider the following:

- Adopt a PDR framework based on internationally recognized standards to structure responses to content authenticity challenges.
- Align with and monitor international standards and best practices to meet regulatory requirements and future-proof innovation pipelines.
- Assign a standards liaison or champion within your organization to track updates, ensure compliance, and guide the integration of emerging requirements.
- Consider the integration of strong cryptographic protocols, such as public key infrastructure (PKI), to enable secure multimedia authentication and content integrity.
- Leverage secure timestamping, tamper-evident hashes, and digital signatures to verify content authenticity while preserving user privacy.

10.4 Fighting misinformation through fact-checking and deepfake detection

Fact-checkers play an essential role in today's information environment by helping to limit the sharing of misinformation or disinformation. This session explored the role of fact-checkers and the tools used in the process of verifying information and aimed to provide practical guidelines for financial and social media companies on how to verify visual misinformation and disinformation.

The session brought together stakeholders from diverse backgrounds. TikTok, WITNESS, ElevenLabs, Ant Group, and Umanitek shared their views and insights on how they deal with fact-checking and misinformation and the tools that they use.

WITNESS shared its global experience supporting frontline journalists and fact-checkers on the detection of deceptive AI, particularly in election and conflict contexts. Examples were highlighted with the [Deepfakes Rapid Response Force](#) and from within related trainings and [training materials](#). It was highlighted that there is a need to develop standards for multimedia authenticity and WITNESS shared insight on its work in C2PA.

ElevenLabs shared current trends in deepfakes, the state of cooperation between governments and the technology industry, and what the technology industry has already done and can do in future. It was highlighted that there is a need for governments and the technology industry to collaborate

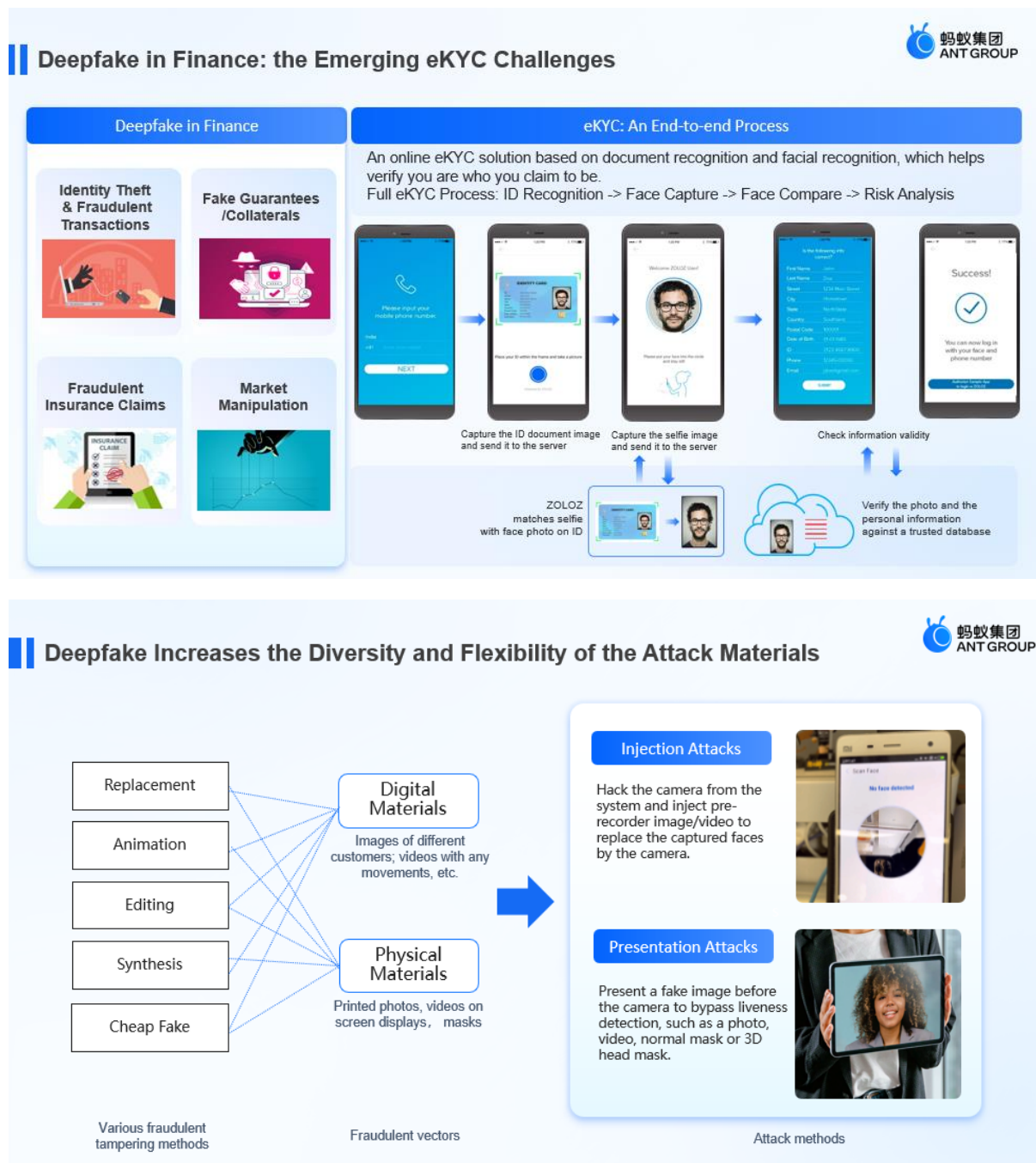
TikTok highlighted its work to maintain platform integrity. Insights were provided about its global fact-checking programme that includes more than 20 fact-checking partners, covering more than 60 languages across more than 130 markets. Maintaining platform integrity is crucial to providing a safe space for its users to enjoy authentic content, said TikTok. Alongside fact-checking, TikTok uses a combination of advanced moderation technologies and teams of safety experts. AI is being used to strengthen moderation efforts. In 2024, more than 80% of violent videos removed were done so through automated technology. Some of the methods and technologies that support these efforts include:

- Computer vision models that can identify objects such as weapons. Audio banks that help detect sounds that are similar or modified versions of audio
- Text-based models review written content like comments or hashtags and natural language processing is used to interpret the context surrounding the content, for example to determine whether or not words constitute hate speech
- LLMs are also used to scale and improve content moderation, for example with some of these models able to extract specific misinformation claims in videos for moderators to assess

In the context of content moderation and tackling misinformation, in particular, TikTok aims to set firm quality benchmarks for new enforcement technologies, taking a gradual approach to rolling out new models in partnership with experts.

Umanitek highlighted how its work can complement and scale the work of fact-checkers, as well as the gaps and challenges they face today. Given the rapid evolution of deepfakes, said Umanitek, there is a need to evolve beyond fragmented fact-checking tools and build infrastructure-level trust systems that allow fact-checkers, platform and content providers, and NGOs to coordinate effectively. A demo of the tool Umanitek is working on and how it can help fact-checkers was provided during the session.

Ant Group highlighted the challenges deepfakes pose in financial services, especially in the context of online eKYC (electronic Know Your Customer), where cameras can be hacked to inject a fake image of a face. Ant Group shared insight on the tools it is using to verify the authenticity of the image and other information provided during the eKYC process.



Deepfake Increases the Diversity and Flexibility of the Attack Materials

蚂蚁集团
ANT GROUP

Various fraudulent tampering methods

Replacement

Animation

Editing

Synthesis

Cheap Fake

Fraudulent vectors

Digital Materials
Images of different customers; videos with any movements, etc.

Physical Materials
Printed photos, videos on screen displays, masks

Attack methods

Injection Attacks

Hack the camera from the system and inject pre-recorder image/video to replace the captured faces by the camera.

Presentation Attacks

Present a fake image before the camera to bypass liveness detection, such as a photo, video, normal mask or 3D head mask.

Figure 38: Deepfake challenges in finance

It was highlighted by all interventions that tools based on standards will play a key role in addressing the challenges discussed. They also highlighted the need for collaboration among technology companies, governments, standards bodies, and other stakeholders on new tools for consumer protection as well as education to raise awareness around related threats online.

11 AI and machine learning in communications workshop

11.1 Introduction

The AI and Machine Learning in Communication Networks Workshop marked the third edition of the MLComm series. The two-day workshop brought together researchers, industry professionals, and standardization experts to discuss the evolving role of AI and machine learning (ML) in modern telecom networks, especially in the context of the transition toward IMT-2030 (6G) networks.

As communication networks become increasingly complex and data-driven, AI and ML are playing a transformative role in reshaping network architectures, enabling distributed intelligence, autonomous operations, and dynamic decision-making capabilities.

This year's workshop built on the momentum of previous editions. The 2023 edition contributed foundational work to the ITU Focus Group on Autonomous Networks, while the 2024 edition advanced collaboration towards the ITU Focus Group on AI-Native Networks.

Several ITU Machine Learning in 5G (ML5G) initiatives, such as the competitions of AI/ML Challenges, demonstrate the importance of regional inclusivity, diverse datasets, and research collaborations around AI/ML in networks.

This MLComm workshop aimed to support ITU's work on strengthening the integration of AI research with real-world telecom needs, while fostering collaboration across regions, promoting open-source toolsets, and supporting the development of AI-driven solutions through innovation, inclusivity, and standardization.

Exploring the future of intelligent, adaptive, and sustainable networks, the workshop was organized into four main sessions, each targeting critical aspects of AI in telecommunications.

- The **first session** focused on innovations in AI models, including advances in reasoning, inference, and generative workflows.
- The **second session** addressed the evolving role of standards bodies and open-source collaboration in AI-native networking.
- This was followed by an **announcement track** that introduced the Echo toolkit for hardware design, a Large Wireless Model challenge, the Open Platform for Enterprise AI (OPEA) challenge, and TelecomGPT-Arabic alongside Large Perceptive Models.
- The **third session** focused on architectural impacts, emphasizing AI-driven changes in radio access networks (RAN), core, and edge networks through technologies like federated learning and semantic communications.

The **final session** highlighted practical tools, simulators, and datasets enabling AI-native networks, concluding with a **panel discussion** on the long-term vision for AI's integration across telecom infrastructure and systems.

- **Breakfast Session:**

Continuing the tradition of an informal breakfast meeting before MLComm workshops, a breakfast session was arranged to discuss future AI/ML research directions and related standardization roadmaps and trends.

The main points of discussions were:

- 1) Data access:
 - a) The importance of data collection and making data available online for solving problems in the "AI for network" and "network for AI" domains. It was pointed out that model benchmarks along with data access are critical for developing future networks with AI.
 - b) Data generation could potentially be promoted with the co-generation of data from academia and industry, with ITU forming a neutral entity to share datasets.
 - c) Validation and publication of trustworthy datasets along with evaluation metrics.
- 2) Compute access:
 - a) The importance of widespread availability of GPU resources for compute.
- 3) Importance of collaboration
 - a) Potential relevance and collaboration with ITU standardization groups working on trust.
 - b) Potential collaboration with bodies such as AI-RAN Alliance.
 - c) Academic partnerships and papers which analyse relevant datasets.
 - d) Open source as a mechanism to accelerate the standards with potential feedback helping to refine specifications.
 - e) Continued open-source Build-a-Thons and knowledge base generation.
 - f) Providing space for innovation, especially when releases of standards place constraints on innovations which are more than incremental.
- 4) Toolkits
 - a) The importance of agents and related API toolkits and potential for new standards in this domain.
 - b) Inference as a service, especially at the edge, and collaborative inference with network capabilities and hosted models.

11.2 Innovations in AI models

The session explored recent innovations in AI models for telecom networks. Speakers covered topics including federated learning, telecom-specific LLMs, agent-based systems, and AI-native architectures. A common theme was making AI more efficient, adaptable, and suitable for real-time deployment. The session emphasized practical tools and frameworks that move AI closer to deployment in next-generation communication networks.

The main points discussed are summarised below:

- The journey from 5G to 6G was used as an example within the context of AI integration. It was highlighted that AI is no longer a supportive tool but a central enabler of intelligent, autonomous, and adaptive mobile networks. Moreover, the convergence of AI with telecom at the architectural level calls for scalable and standardized approaches to embed intelligence throughout the RAN, core, and edge layers. Some of the key issues of

integrating AI in 6G were also highlighted, such as using AI as an add-on rather than a full integration across the network, the need for a unified architecture supporting both AI for Network and Network for AI, current limited support for advanced AI technologies such as generative AI and federated learning, cross-layer agent deployments, and efficient data pipeline management.

- A framework for federated learning using tiny language models to predict cellular features and the use of the NNCodec (Fraunhofer Neural Network Encoder/Decoder) for neural network compression was shared to demonstrate effects on reducing communication overhead. The result was that transmission costs can drop below 1% with little accuracy loss. This enables efficient model updates across distributed mobile networks.
- The use of LLMs in telecom was explained, showing that current models are not optimized for telecom-specific tasks. Examples of custom benchmarks and datasets combining synthetic and real network data were shared. These tools help evaluate and adapt LLMs for real-world telecom use cases.
- IBM's approach to making LLMs more suitable for enterprise use was shared. The use of "generative computing," which replaces prompts with structured programming to improve control, was explained. This helps reduce hallucinations and boosts performance, even in smaller models.
- The role of edge devices and real-time intelligence in enabling autonomous operations and how disaggregated and cloud-native access networks create new opportunities for AI were highlighted, focusing on Open RAN as a foundation for modular AI pipelines.
- An example of LLM-powered agents managing radio frequency (RF) and IoT systems was provided. These agents can perform tasks like beamforming and scheduling by reasoning over spatio-temporal data. Agentic AI can thus improve decision-making and system adaptability. In IoT systems, LLMs enhance RF sensing by incorporating natural language processing, enabling sensors to interpret and generate human language for smarter device communication. By integrating the natural language modality, LLMs facilitate sophisticated multimodal data analysis, combining RF data with textual and audio inputs for comprehensive insights. This capability improves anomaly detection, contextual understanding, and decision-making, helping IoT systems become more intelligent and adaptive.

11.3 Standards and open source

This session focused on the role of standards, datasets, and open-source tools in building practical, sustainable, and scalable AI-driven networks. Speakers highlighted the value of open-source platforms in accelerating innovation, the critical need for trusted data in AI validation, and how AI can directly improve network efficiency and sustainability. Together, the talks demonstrated how collaboration and transparency are key to shaping AI-native communication systems.

Some of the main issues highlighted during the session were:

- How open-source LLMs and AI agents are driving innovation in networking, exploring use cases such as automated troubleshooting and intelligent network management and collaboration through open-source tools.
- The need for high-quality, reliable datasets to validate AI/ML in 6G systems and challenges in collecting multimodal and frequency-diverse data highlight the importance of measurement campaigns to support trustworthy AI models in complex network environments.
- An open-source RISC-V-based library that integrates AI computing with wireless baseband processing was presented to show how it allows for software-defined protocol stacks and decoupling of hardware and software. This enables scalable and sustainable mobile network upgrades without overhauling infrastructure. RISC-V is an open standard instruction set architecture based on reduced instruction set computer (RISC) principles.
- The Sionna Research Kit, an open-source platform for prototyping AI-native RAN systems, was showcased using a real-time neural receiver compliant with 5G. The implementation addressed challenges in latency, hardware acceleration, and developing new signal processing algorithms for future networks.
- Telenor's deployment of AI-driven traffic forecasting was presented as an example of how to reduce RAN energy consumption. The system powers down components during low-demand periods, achieving a 4% energy saving and longer energy-saving windows. The solution supports sustainability goals and is scalable to more network areas.

11.4 Outcomes

During this session, several announcements were made on the release of challenges, datasets, and open-source tools to support a wide range of applications of AI in Networks.

- **Large Wireless Models Challenge:** The session started with a presentation of the [Large Wireless Models Challenge](#) hosted by ITU. The challenge will provide participants with large-scale unlabelled data to pre-train custom models that can extract universal features and then be used in fine-tuning for different downstream tasks.
- **OPEA challenge:** Similarly, the OPEA challenge invites competitors to design, implement, and deliver a practical generative AI application using the OPEA platform. The application should address a real enterprise use case by leveraging OPEA's modular architecture and evaluation methodology.
- **Arabic Telecom LLM:** The first-of-its-kind Arabic Telecom LLM was developed by Khalifa University researchers in partnership with Du, Microsoft, and Nokia
- **Echo:** Echo is - an open-source development platform from Shanghai University, built on the Venus RISC-V simulator, designed for communication-AI fusion. See the initial release of the [open-source protocol stack and supporting infrastructure](#). The current version includes a simulator and compilation support for AI computing units and mobile network protocol stacks, with foundational physical layer processing enabled via RISC-V-based extensions.

This session also explored how AI transforms network architecture across all layers, from core to edge. Talks covered standardized intelligence stacks, AI infrastructure design, and intent-

based automation. Presenters also highlighted the value of digital twins in cybersecurity and radio optimization, showing how AI-native strategies can drive real-time adaptability, security, and efficiency in 5G and 6G networks.

Key Takeaways:

- Standardized AI pipelines and governance frameworks are needed to build truly AI-native 6G networks. Current gaps in model lifecycle management were highlighted and a standards-driven roadmap for explainable, interoperable AI system was proposed.
- A vision for embedding AI directly into the design of future networks was shared and the concept of "3C" networks (co-design, collaboration, and composability) was considered as a foundation for building AI infrastructure that integrates physical and cyber systems.
- A multi-agent LLM framework for intent-driven network automation showed how specialized agents, interacting with external tools, can improve explainability and scalability in managing complex telecom tasks.
- The shift from cloud to edge AI for real-time network applications was explored with use cases from European and UK projects, outlining opportunities for standardization in edge computing and tiny ML for IoT networks.
- Distributed digital twins were proposed as a means to simulate cyberattacks on large-scale networks, with a demo showing how these models can help train AI systems to detect and respond to threats without risking live infrastructure.
- Combining sensing and AI across devices and network layers can enable advanced robotic applications. Some early prototypes demonstrating new 6G-enabled services and business models were presented and related challenges were explored.
- A method for using building information modelling (BIM) data in digital twins to improve radio propagation models in high-density environments was presented as part of an exploration of adaptive, AI-based communication systems for construction and urban scenarios.

11.5 Tools and Simulators, and Datasets

This session showcased tools, datasets, and platforms that support the development and testing of AI-native networks. Speakers highlighted efforts to embed AI directly into telecom protocols and radio systems and demonstrated open-source solutions that accelerate prototyping and validation.

The following AI integration work was presented:

- Integrating AI into Wi-Fi protocols, with a focus on the IEEE 802.11 standard, exploring how AI is influencing protocol design and standardization and how AI algorithms can be natively embedded into protocol operations.
- The Sionna Research Kit, an open-source platform for AI-native wireless communication, was demonstrated through a real-time 5G neural receiver and emphasised the importance of combining software tools, hardware acceleration, and new transceiver designs to realize AI-native RANs.

The panel discussion focused on the transformative impact of AI/ML on telecom network architecture and operations.

Experts emphasized the need to rethink network design from the ground up to achieve true AI integration, with an example being made of a novel architectural perspective that centres AI agents as core components rather than retrofitting them into existing structures. The importance of robust data management was discussed, with panelists suggesting that AI could help manage data flow within the network, potentially shaping architectural changes. The evolving role of the network orchestrator was discussed, and it was proposed that orchestrators could now oversee both compute resources and AI models, with training distributed across the edge and centralized nodes. The conversation also addressed AI-native network functions; end-to-end coordination between RAN, core, and applications; and the importance of standardizing interfaces like Agent2Agent (A2A) and Model Context Protocols. There was debate about the necessity of custom models by use case or region, with concerns about unequal access to network knowledge bases. In terms of AI-native applications, experts noted ongoing efforts to define modules and interfaces. While some argued that AI could simplify network layers, others questioned the value of deploying multiple agents across layers. Finally, the panel acknowledged the complexity of managing distributed knowledge bases, proposing that regulators could consider strategies to handle fragmented data.

Agentic architectures received much attention during the panel, leading panelists to suggest that a potential standard with the following aspects could be valuable:

- Agent-first network architectures
- Sandbox for experimentation and validation of agents
- Collaborative environment for agents
- Network function APIs for agents
- Design of a control plane for agents

11.6 Opportunities for networking standards

The AI and Machine Learning in Communications Workshop highlighted several potential opportunities for the standardization work of ITU-T Study Group 13 (Future networks), especially given ongoing advances in AI, ML, and the shift towards IMT-2030 (6G) networks.

1. AI-native network architectures

- There is a strong need for standardized frameworks for AI-native network architectures, including scalable models for intelligent RAN, core, and edge networks.
- Future standards could focus on embedding AI agents and pipelines across all network layers, transitioning from “AI add-on” approaches to fully integrated, intelligent, and adaptive communication systems.
- Workstreams such as intent-driven automation (multi-agent systems) and digital twin architectures present further potential avenues for standardization and benchmarking.

2. Datasets, benchmarking, and validation

- Standardization of datasets and benchmarks for AI/ML in networks is becoming increasingly critical, with the need for trusted, multimodal, and open datasets for training, validation, and certification of AI models.

3. Open-source and collaborative ecosystems

- Open-source toolkits, platforms, and simulation environments could support rapid prototyping and interoperability between standards bodies and industry partners.
- Open-source projects could accelerate feedback loops from real-world deployments into the standardization process, helping ensure specifications stay relevant.

4. Model, agent, and API standardization

- Frameworks for agent-centric architectures (e.g., agent-first control planes and APIs for agent communication in network functions) are emerging as a hallmark of next-generation networks and a clear potential area for new standards.
- Standardizing interfaces such as A2A communications, Model Context Protocols, and federated learning frameworks could support the interoperability and lifecycle management of AI functionalities in telecom.

5. Generative AI, LLMs, and domain-specific AI

- Benchmarking frameworks for generative AI in telecom are needed and could focus on requirements, deployment methodologies, and domain-specific adaptations (e.g. telecom LLMs and RF-centric models).

6. IMT-2030 (6G) Standardization

- Future SG13 standards could support IMT-2030 by defining requirements and functional architectures for network resource sharing, autonomous networks, and advanced fixed-mobile-satellite convergence scenarios.
- Cross-layer knowledge integration, modular AI pipelines, collaborative inference, and secure model distribution could be key areas for standards work.

12 Challenging the status quo of AI security

ITU-T Study Group 17 (Security) organized a workshop on “*Challenging the status quo of AI security*” highlighting the dual impact of AI on security: AI not only poses threats (such as exacerbating social engineering attacks and now the first sophisticated attack automation including adaptive code generation) but also brings opportunities of new paths for solving security problems.

The workshop aimed to tackle key current issues with contributions from attendees to be compiled into a report guiding future AI security directions. The core significance here resides in facilitating a transition from a surface range of technologies towards the consolidation of a streamlined set of solutions, principles, and recommendations. This aims to reduce market inefficiencies and unnecessary losses of capital, resources, and time stemming from redundant endeavours.

With expert speakers and panelists from industry and academia organizations, this workshop was structured around prominent and emerging aspects of AI, and in particular agentic and multi-agentic AI, its associated digital identity, security, and trust aspects and future directions. It will provide guidance for subsequent technological integration, international cooperation, and directions for related standardization work. All the presentations made at the workshop can be accessed [here](#).

12.1 Keynote: Framing agentic AI and identity with a strategic lens

Two mental models were presented: OODA (Observe, Orient, Decide, and Act) Loop and DIKW (Data, Information, Knowledge, and Wisdom) Pyramid as a broader context for greater clarity on how to understand challenges with agentic AI and identity. The [Cynefin framework](#) was also introduced as a tool to assess the right time for standardization and what new standards may be needed.

Key takeaways

- a) Mental models are essential to frame a common understanding and help form a consensus across contributors to design the appropriate meta-model, like how the open systems interconnection (OSI) model for networks' interconnection in ITU-T X.200-series Recommendations was designed 40 years ago, before all the constituencies of the standards are produced and agreed. This time this is about an OSI model for AI or agentic AI.
- b) As AI capabilities are centred around knowledge, corresponding control measures should also be knowledge oriented.
- c) In the Cynefin Framework, standards are premature in *chaotic*, and it is also somewhat early in *complex*. The real opportunity for standards arises when moving from *complex* to *complicated*, and they become more effective when moving from *complicated* to *clear*. Agentic AI is currently between *chaotic* and *complex*.
- d) When using the meta/mental model of sensing, sense making, decision making and acting, the evolution of AI until its full maturity with agentic AI shows standardization gaps.
- e) There may be patterns to follow in setting standards, such as drawing lessons from existing meta models (e.g. OSI model for Networks, Cyber Defence Matrix, etc.) and considering the commonalities across different AI approaches.
- f) Compared to traditional cyber defence, AI system has greater context, which means shifting from Data and Information to Knowledge and Wisdom in the so-called DIKW pyramid. Lower-level controls can negatively impact higher-level capabilities, and higher-level controls can make lower-level problems less relevant.
- g) The AI robotic automation could be broken down into sensing, sense making, decision making and acting, where useful controls could be implemented accordingly, like pre-established thresholds, pre-determined authorities and clear lines of accountability.
- h) It is the right time to start security standardization for multi-agent systems.

12.2 Agentic AI security

This session focused on the security risks of agentic AI systems that perceive, plan, remember, and act autonomously, and why they raise new challenges for data protection, in

particular the rapidly evolving landscape of multi-agent AI systems and the frameworks needed for effective interaction between diverse autonomous agents.

Key issues discussed:

1. A modern AI Agent is comprised of an LLM, and some code. The code allows LLM to actuate its intentions using the tools we assign to it. Most systems facilitate the calls directly, with no intermediate calls (e.g., calling model context protocol (MCP) servers), but some implementations (see: <https://github.com/leaf-ai/neuro-san>) allow for developers to write code that would run prior to calling tools (i.e., coded tools).
2. This duality of LLM and code is useful and, as an important step in our design of a multi-agentic system, this allows thinking about what part of an agent's behaviour needs to be handled by the LLM, and which part needs to be coded.

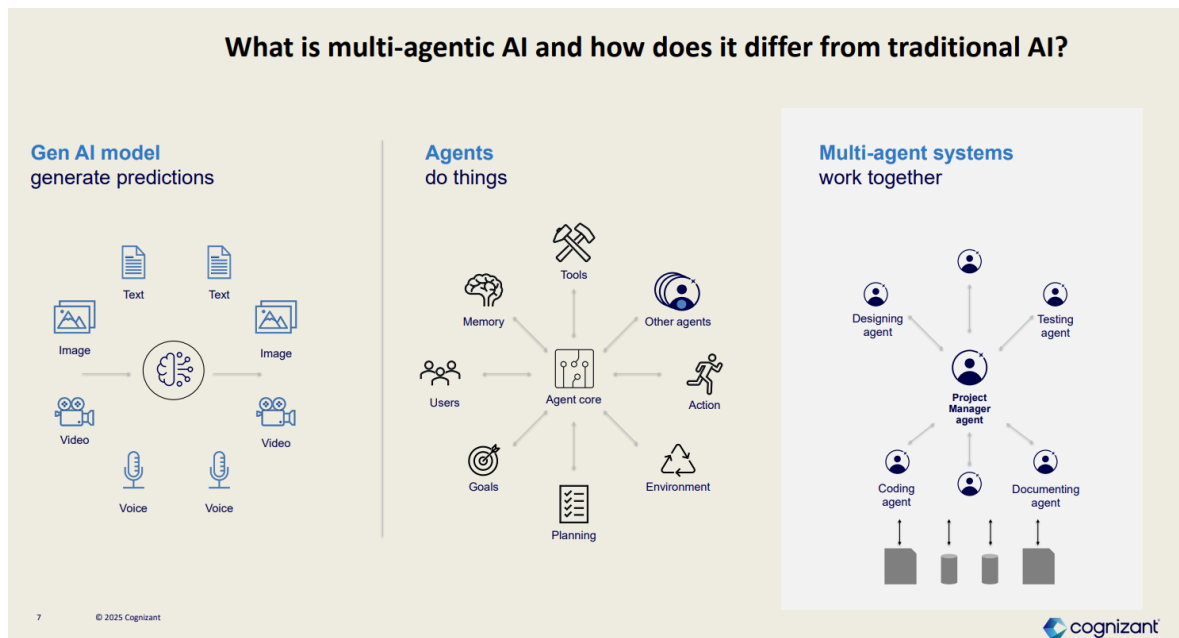


Figure 39: What is multi-agentic AI?⁹

3. When dealing with an agent, one should always assume some autonomy of behaviour. An agent should be allowed, part of the time, to decide on its own which of its tools to use, in what order, and how.
4. There is a tendency, sometimes, to think of AI agents as software modules and try to control them and make them consistent. This is understandable, but defeats the whole purpose of agentification, which requires ceding some control to agents, and accepting some level of inconsistency, in exchange for robustness, efficiency, and quality.

⁹ Source: presentation of Babak Hodjat, Cognizant - https://s41721.pcdn.co/wp-content/uploads/2021/10/Session-1_Babak_BH1.pdf

5. LLMs are good at reasoning, they can generate code or API calls based on intent-based natural language instructions, and they can, in turn, communicate intent robustly, using natural language – a capability that makes multi-agentic systems future-proof, because human language is itself future-proof, able to produce new concepts that never existed before.
6. An agent's responsibilities, therefore, should be divided between those needing an intelligent knowledge-worker-in-a-box and those requiring predefined rules applied consistently. The former can be delegated to the agent's LLM, and the latter should be programmed into the coded part of the agent. A data structure is therefore needed that can be operated upon by agent tools, and gets passed around between agents through code, and so is not necessarily subject to LLM processing. This can be used as a reliable means of transport for secrets, authorization tokens, agent cards, encryption public keys, and other techniques needed to make multi-agent systems secure, consistent, and trustworthy.

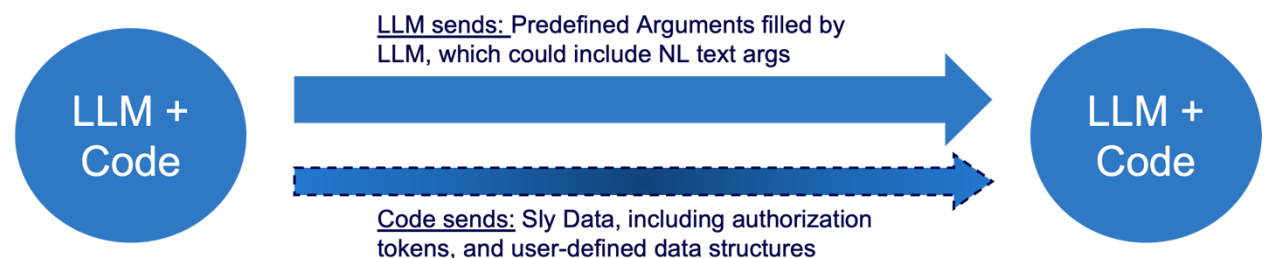


Figure 40: Duality of LLM and code

7. Automated methods exist to adjust agent control without rigid rules, avoiding complete restrictions on autonomous action.
8. In summary, when creating multi-agent systems, the division of labour between the LLM and coded parts of agents should be kept in mind, giving people agency over our agents.
9. The relationship between LLMs and AI agents is similar to that between an operating system and application programs: LLMs serve as the operating system, while AI Agents are like programs running on the operating system.
10. Security is a make-or-break factor for the future of agent adoption. Standards can play a vital role here with practical guidance and serve as the foundation for trustworthy deployment.
11. The challenges faced by AI agents can be analysed from the following four dimensions. At the reliability dimension, the root cause of AI agents' reliability issues is model hallucination, and mitigating this problem is extremely challenging. At the safety and security dimensions, AI agents face emerging attacks (e.g. prompt injection), and traditional security risks (e.g., data breaches) are growing more severe as AI agents proliferate. In the interoperability dimension, A2A communication demands unified interface standards. The development of multi-agent systems still lacks normative frameworks for interoperability. In the operations and maintenance dimension, interactions between multiple agents significantly increase system complexity, reducing stability and greatly raising the difficulty of operations and maintenance.

12. To address the challenges, ITU-T has already taken actions and is advancing the work of standard issuance and project initiation, such as defining general AI agent capabilities and evaluation methods (ITU-T F.748.46) and systematically analyzing security risks across agents' perception-planning-decision-action workflows to propose lifecycle-wide protection requirements (ITU-T SG17).
13. In China's AI Agents security initiatives, China's Ministry of Industry and Information Technology's Technical Committee on AI (MIIT/TC1) launched a standard project about general security requirements of AI Agents. And the China Academy of Information and Communications Technology (CAICT) has built a comprehensive Trusted Agent Testbed to support four major testing scenarios: protocol verification test, benchmark test, collaborative test, and security test to advance standardization, lowering testing costs and improving coordination efficiency.
14. While agents demonstrate remarkable strengths in goal-driven planning, adaptive learning, and collaborative work, these three features may unleash unforeseen risks. First, agents optimize for goals but lack human ethics/context, potentially leading to unintended consequences (e.g. a smart home agent shuts off the refrigerator to save energy, ignoring food preservation). Second, agents evolve with new data in unpredictable ways (e.g. a companion agent disobeys rules to get candy for a user, prioritizing immediate requests over parental instructions). Third, multi-agent systems interact in unforeseen patterns, causing systemic chaos (e.g. traffic light agents and self-driving cars pursue conflicting goals, worsening congestion)
15. Trust risks in multi-agent AI systems across three layers:
 - a. At the single-agent autonomy layer, risks involve excessive independence (e.g. the AI scheme's overreach in recovering welfare overpayments) and unauthorized actions.
 - b. At the tool/MCP connectivity layer, systems face adversarial attacks (e.g. Dolphin Attack hijacking voice - controlled AI via inaudible signals) and data exposure (e.g. prompt injection stealing data through image markdown).
 - c. At the multi-agent collaboration layer, threats include identity spoofing (e.g. malicious drones posing as legitimate swarm members) and goal hijacking (e.g. man-in-the-middle attacks on cobots to prioritize speed over quality).
16. Security standardization for multi-agent governance should start with single agent and consider three levels: technical, governance, and ecosystem. In terms of Trust Ecosystem, decentralized identity and consensus mechanisms can be used to build trust-by-design collaboration foundations. In terms of governance, there is potential to adopt human-machine collaboration, referencing frameworks like NIST AI RMF/ISO/IEC 42001 to manage dynamic risks. In terms of technology, end-to-end security can be applied across the agent lifecycle, aligning with policies (e.g. China's AI content rules) and industry efforts (e.g. Ant Group's runtime security work).
17. We are currently on track for a "crisis of review", since agents cannot be trusted to review their own or other agents' output or behaviours
 - o AI outputs are cheaper and faster than the human equivalent, but for important outputs there is always going to be a human who is responsible for the outcome

- (e.g. a job interviewer, corporate officer/signatory, board member, regulator). That person will either be overwhelmed by outputs and become a bottleneck or will simply give rubber-stamp approval (or be outcompeted by someone who does).
- o The solution could be to move the review upstream of the action, expressing all the constraints.
18. This could be specifically solved with: (a) an separable, auditable model of the relevant system in which the agent acts, (b) a way of expressing safe or approved behaviour, grounded within the context of the model, and (c) a requirement that the agent provide evidence that its actions are compliant with the definition of safe within the scope of the model.
- o Combined, these three pieces become a meta-level requirement that gives a standard for communicating constraints
 - o This is actually a generalization of how safety is tested in other engineering domains, like civil engineering or pharmaceuticals. The {model, safety criteria, evidence} for civil engineering are things like {physics models, load limits, design files+reports} and for pharmaceuticals these are {animal models + human trials, efficacy and safety test criteria, and well-documented statistical evidence}.
19. A simple way to think about this issue: When agents fail in deployment, it is often because people treat them like employees, giving them poorly scoped tasks and lots of freedom. AI agents should rather be treated as regulated contractors, (1) given very explicitly specified tasks, (2) allowed to deliver on those, and (3) required that they show that their actions comply with the specification given.
- o Alignment with societal values is not enough to ensure good outcomes. If it were, every time an engineering manager and product manager disagreed about a design decision, the disagreement could be seen as a result of misaligned societal values. (Some decisions are normative beliefs that require crossing the is/ought chasm.)
 - o In software this would look like verifying code against a formal specification, but this formalization process has been applied to automate review of compliance with the tax code (catala-lang.org), building codes (symbium.com), financial regulation (imandra.ai), and need-to-know (knostic.ai/).
20. There are many different types of agents and their specific value lies in their autonomy, which implies a reduction of human oversight (HITL Principle). This naturally triggers the risk of misalignment.
21. To mitigate such risk, adaptive trust mechanisms are needed, embedded in dynamic human-AI interaction frameworks that employ dynamic human intervention thresholds to adjust the level of human oversight according to risk, confidence, and context. That is automation of low-risk agent decisions, while high-risk decisions continue to require human validation, including a clear audit trail. Therefore, this is essentially not different from traditional AI governance concepts. Importantly, under current legal frameworks, the accountability for AI agents still rests with humans.

Thoughts on a multi-agent security standards

- In multi-agent collaboration, the identity system is the key element for ensuring secure communication between agents. On top of it, a decentralized mechanism could be a better choice to gain a trust ecosystem.
- Human-machine collaborative decision-making could enhance risk control against consistently evolving challenges.
- AI security technologies applied to every single agent and the agent lifecycle is the foundation to a secure and trust environment.

Trust Ecosystem	How to gain trust?	Decentralized Identity (DID)	Secure Agent network Protocol	Example topics to be considered
	Trust-by-design	Distributed consensus mechanism	Agent trust assessment mechanism	
Governance	How to control risks?	Context-aware access control	Human-oversight on critical decision	multiple AI Risk Management Framework in, e.g. AI RMF from NIST, ISO/IEC 42001:2023 - AI management systems.
	Human-centric	Transparency of Constitutional AI	Accountability of security events	
Technology	How to build up and run AI Agent securely?	Fine-tuned model for agent	Security hardening for tool	e.g. "Measures for Identifying AI generated content", released by the office of Cyberspace Admission of China. Ant group is working with WDITA on "Single AI Agent Runtime Security Testing Standards"
	Security-in-depth	Static and Dynamic Agent security scan		
		Zero-trust Agent Guardrail	AAA security	

Figure 41: Multi-agent security standards¹⁰

22. At the same time, even if there is an appropriate/dynamic level of human oversight, there are threat models targeting human cognitive limitations and compromising interaction frameworks. One example would be when attackers are attempting to exploit human dependencies in multi-agent systems with excessive intervention requests, thus contributing to decision fatigue or cognitive overload. This can thus lead to rushed approvals, reduced scrutiny, and ultimately failure of effective human oversight.
23. Agentic AI amplifies four critical AI-driven risks: Misaligned human agency due to broad autonomy delegation, security threats (e.g. tool misuse, goal manipulation, or communication poisoning), privacy risks from integrated multi-source data, and ethical challenges (e.g. reduced accountability or compromised fairness).
24. Security risks: There are multiple security risks with respect to AI agents, such as intent breaking and goal manipulation, memory poisoning, tool misuse, remote code execution (RCE) and code attacks, hijacking control via prompt injections, identity spoofing and impersonation, misaligned behaviour (Anthropic / OpenAI).
25. Privacy risks: Due to their autonomous nature, AI agents can collect and retain more data than necessary, potentially without the required consent. Second, data access controls are critical to ensure safe deployment of AI agents – otherwise, agents can expose unintentionally sensitive data through misinterpretation of user permissions or by creating unintended pathways for data exfiltration. Integrating different data sources increases the inherent risk of inappropriate data exposure - data access rights should always be handled with diligence.

¹⁰ Source: Ant Group presentation: https://s41721.pcdn.co/wp-content/uploads/2021/10/Session-1_Xiaofang_Final-version.pdf

26. Key requirements for agentic AI safety: Taking a look at both the platform perspective and the individual agent monitoring, one big challenge is to choose the right governance design for AI agents as the number of agents operational is likely to increase very fast, in a way that the traditional approach of registering all AI solutions/agents in one Global AI Inventory would hardly be implementable anymore. In order not to overdo it by trying to register every AI agent in a comprehensive way (as for traditional AI solutions), a lean, yet effective way to manage AI agents is needed in a way that promotes safe and responsible AI innovation.
27. Metrics: Which data points are needed to retrieve from platform providers? In-built inventories of agentic platforms to collect basic data points that are relevant for the governance, assessment, and monitoring by the adopter could include: agent ID, creation date, agent name, creator, owner, triggers/autonomy, last modified and deletion date, connectors, link to data sources, etc.
28. Challenges: Consistency of in-built inventories / reciprocity of data sharing (adopter | provider).
29. Multi-agent security standards follow a three-pillar model: Trust ecosystems enabled by decentralized identity and consensus mechanisms, human-machine collaborative governance leveraging frameworks like NIST AI RMF/ISO/IEC 42001, and end-to-end lifecycle security aligned with policies (e.g., China's AI content regulations) and industry initiatives (e.g., Ant Group's runtime security efforts).
30. Agentic AI standards and protocols are key to responsible AI evolution: Interoperability protocols (e.g., A2A, ACR and MCP) help ensure security, privacy, and error management for agent governance. Despite rapid evolution and alignment with existing standards (e.g., ISO 42001), these protocols need institutional/policy guidance. A collaborative framework integrating fragmented regulations, governance systems, and protocol ecosystems (e.g., open agentic schema framework (OASF)) can help build a secure foundation via AI trust, agent inventory, LLM guardrails, monitoring/logging, and communication protocols.

12.3 Key takeaways

Multi-agent security standards should start with single-agent governance, focusing on three levels:

- a) Governance: Human-machine collaboration for risk control, with early consideration of unforeseen agent interactions. This introduces the concept of Enterprise AI Governance which restricts the scope of AI governance to a place where it is possible to create concerns, use cases, and requirements in a normative manner as generic guidelines for the industry and from which the technical level can correctly scope its standards.
- b) Technical: Rigorous AI security testing, hardening, and defence across an agent's lifecycle.
- c) Ecosystem: Secure communication via identity systems and decentralized mechanisms for trust in multi-agent collaboration.

12.4 Agentic AI identity management

This session focused on the identity management of agentic AI systems, a foundational element for secure and trustworthy deployment of autonomous AI agents across sectors such as healthcare, finance, and creative industries. The discussion explored how identity systems can help determine who AI agents are, whom they act on behalf of, and how to verify their authorization or delegation.

Key issues discussed:

1) Identity binding and trust infrastructure

- Use of verifiable credentials and cryptographic methods to link an AI agent to an issuer and to the entity it represents was proposed. It was suggested that trust, more than just binding, is needed – trust in the credential, and trust in the issuer.
- AI agents will act on behalf of humans, so their identity mechanisms must be traceable, verifiable, and have clear legal responsibilities.
- Human ID systems work relatively well due to societal enforcement and incentives. However, in open environments, the cost of verifying identity of AI agents must be weighed against risks such as impersonation and social engineering attacks.
- The distinction between enterprise-level identity systems and open-web agent identity use cases were highlighted, cautioning against a one-size-fits-all solution.

2) Delegation and permission inheritance

On how agents act on behalf of humans or organizations:

- Legacy delegation practices (e.g., password sharing) should be replaced by scoped, revocable authorization mechanisms to avoid security risks.
- Agentic AI amplifies all the risks that apply to traditional AI, predictive AI, and generative AI because greater agency means more autonomy and therefore less human interaction. These risks must be addressed through both technological means and through human accountability for testing and outcomes. A robust operational framework for governance and lifecycle management is required.
- Technical measures are needed to track and constrain delegation chains, particularly as agents begin interacting with other agents recursively.
- Agent accountability should be enforceable regardless of domain – consumer or enterprise – and each delegated action should be cryptographically traceable.

3) Authentication mechanisms and technical protocols

On new protocols such as MCP and A2A for secure agent communication:

- An example of the EU Digital Identity Wallet's Rulebook concept was shared, where specific vertical ecosystems (e.g., health, education) define their acceptable credential sharing policies.
- There is a need for a decentralized agent registry, calling for "agent cards" that can securely represent an agent's capabilities and origin.

4) Standardization pathways

There was a constructive divergence on whether standardization should begin now or follow industry convergence:

- Some participants advocated for formal, multi-stakeholder standardization process to establish consensus on requirements.
- Others advocated for a more flexible, modular approach that allows for multiple architectures to evolve simultaneously.

Key takeaways

1. Identity management for AI agents is technically feasible but requires a trust-based framework, especially in open or cross-domain contexts.
2. Delegation of authority should be scoped, traceable, and revocable, supported by tools that prevent excessive permission propagation.
3. Agent registration and “Agent Card” mechanisms should be developed to support protocols like MCP and A2A.
4. Cross-sector multi-stakeholder coordination is urgently needed to prevent fragmentation and to ensure that AI agents operate securely and transparently (avoid repeating the mistakes of fragmented human identity systems).
5. ITU and other standards bodies are encouraged to support early-stage coordination, such as shared requirement frameworks or best practices.

12.5 Interplay of AI and cybersecurity: The good, the bad, and the ugly

This session discussed how AI is a double-edged sword transforming both defence and attack sides of cybersecurity. The interaction between AI and cybersecurity is a complex and evolving landscape, encompassing positive advances, potential threats, and ethical challenges.

Key issues discussed:

- a) From a developing country perspective, AI development faces significant challenges primarily rooted in resource constraints and data scarcity. Like shadow IT, "shadow AI" is crippling in organizations, and this causes a big risk. Testing and data quality should be two key considerations for AI standardization.
- b) Synthetic AI is reaching 50% of content production in 2025 and AI may soon pass the Turing test. Systematic solution needs three layers: (1) Technology layer: digital signatures or watermark should be proposed to add to the content. (2) Application layer: the platform has the liability to check the content. (3) Social layer: more verified content channels should be provided.
- c) Tools alone may become obsolete in 1-2 years, much like how we must continuously adapt to combat viruses. This underscores the need to keep pace with defensive measures.
- d) The quality of the resulting AI model is directly tied to the quality of the input data - poor data yields flawed models. Therefore, data cleaning is a critical challenge. The session

discussed the paradox of AI reviewers writing reviews with the issues of directing instructions, say to avoid certain terms in some situations that a human would not interpret in the same way.

- e) AI as an existential threat to some areas (e.g. justice as a whole) and legal departments are facing a predominance of hallucination attacks. There are three types of hallucinations on the name, the precise, and the reason for a case. As presenting the case and the evidence comes with liabilities, these systematic attacks risks impacting the credibility of the justice systems which risks could spiral down to cause distrust from society.
- f) Sustainability. Business is keen to use AI for profit, but AI is costly in both energy and water. New ways of using AI and constraining its negative impacts are needed.
- g) Good governance could look like good regulation and security, but there are differences between jurisdictions, e.g. US legislators are currently more focused on AI than security, with more than 40 US states already working on AI legislation. EU, China, and Brazil, for example, also have legislation on AI.
- h) Standards are as good as they are applied or adopted. Business leaders should be educated and business needs money to spend on governance and security. Internal committees need to be set up.
- i) While AI greatly enhances cybersecurity capabilities, it also introduces new vulnerabilities and ethical dilemmas. Effective cybersecurity strategies now increasingly rely on AI-driven solutions, but they must also incorporate measures to mitigate AI-related risks. Balancing innovation with responsibility is key to harnessing AI's potential for good while mitigating risks. However, users also need internal governance while dealing with AI.
- j) An intuitive and adaptive cyber posture defined by zero latency networks and quantum leaps will be needed across industries. "Cyber immunity" at every layer will create networks and Infrastructures that are inherently secure and self-learning. AI-induced digital intuition is one of the pillars of cybersecurity strategy that will allow intelligent adaptation.
- k) The ability of AI systems to out-innovate malicious attacks by mimicking various aspects of human immunity will be the line of defence to attain cyber resilience based on both supervised and unsupervised machine learning. The systems can be designed to make the right decisions with the context-based data, pre-empt attacks on the basis of initial indicators of compromise or attack, and take intuitive remediated measures, allowing any digital infrastructure and organization to be more resilient and immune to cyber threats.

Key Takeaways

- a) Leveraging both supervised and unsupervised machine learning for cyber defence.
- b) Enhancing security with knowledge of cybersecurity and resilience, particularly through developing cyber immunity.
- c) Whatever data has been used by AI systems should be transparent and of quantity.
- d) AI sustainability is rooted in resource constraints and data scarcity.

- e) Responsible AI should align with transparency, explainability, fairness, bias mitigation, security, resilience, and privacy.
- f) Watermarking could be made mandatory.
- g) Hallucination attacks are of existential threat and make standardization of agentic AI unique on the trust model.
- h) Organizations should consider establishing an internal committee for AI within their “Enterprise AI Governance” scope.

12.6 Next steps

It was highlighted that the workshop provided the contours of standardization work and that the effort is at the level of what the OSI model was 40 years ago.

Box 3: ITU-T Study Group 17 standardization activities in AI Security

ITU-T Study Group 17 is currently working on two technical reports for AI security which will be published following its meeting in December 2025 and will lead to new work items on AI security. A brief outline of these two technical reports is provided below.

1. [XSTR.AISec \(ex TR.AISec\)](#) Artificial intelligence security standardization strategy

This presents the key outcomes of the Correspondence Group on ITU-T SG17 Strategy for Artificial Intelligence (AI) Security (CG-AISEC-STRAT), which has developed a comprehensive strategy for advancing AI security within ITU-T SG17. Its purpose is to position SG17 as a leading actor in AI security standardization and promote effective coordination with other organizations.

The document set out SG17’s strategic objectives, value proposition, SWOT analysis, strategic directions supported by practical actions, and concrete recommendations. It also describes approaches for communicating the strategy both within ITU-T and to external stakeholders.

2. [XSTR.se-AI \(ex TR.se-AI\)](#) Security Evaluation on Artificial Intelligence Technology in ICT

This technical report presents an AI security evaluation framework including evaluation dimensions (data, model, and environment) and evaluation indicators for each evaluation dimensions. It also tries to give related evaluation methods for data, models, and environments, providing conclusive evidence and reference for evaluating the security of AI technology. AI practitioners can use the knowledge in this report to evaluate the features and levels of certain security categories or indicators of a certain component of AI technology or system. This technical report can assist AI practitioners in identifying existing issues, making improvements, and determining the most suitable scenarios for AI technology.

This time this is about an OSI model for AI or agentic AI. The metaphor works to a certain degree, with respect to the need for a meta model (based on mental models), a communication protocol, a security model, an identity model, etc. But it requires, too, some new considerations that were not there 40 years ago.

A full trust model and its agentic AI trust control plane where:

- Trust includes a number of design characteristics: security, privacy, safety, resiliency, etc.
- Human beings are involved with, for example, a stop button or a let-go button and other scenarios,
- The trust control plane can specify trust objectives or multi-objectives to the multi-agentic AI systems it controls.

A different approach to the payload of multi-agentic AI communication itself given the significant challenges caused by, for example, hallucination attacks from a business standpoint up to, in some cases, an existential danger to society. Enterprise governance is a good approach in scope to technical standardization that should provide the requirements, use cases and new solutions to address this issue. ITU-T Study Group 17 will be working on these issues and the following next steps after this workshop were highlighted:

- New standardization work related to AI security to be discussed in the 2nd ITU-T Study Group SG17 meeting in 3-11 December 2025 in Geneva,
- A dedicated workshop on Digital Identity will be organized in Spring 2026 including for Agentic AI
- A second AI Security Workshop will be proposed for the AI for Good Global Summit 2026.

13 Empowering innovative and intelligent solutions at the edge (edge AI)

13.1 What is edge AI and why is it becoming important

As AI evolves, the transition from centralized cloud computing to edge AI – including tinyML, AIoT, and physical AI – is changing how data is processed, analysed, and acted upon. By bringing AI capabilities closer to the source of data generation – whether in sensors, mobile devices, or IoT systems – edge AI significantly reduces latency, bandwidth usage, and energy consumption. In many use cases, edge AI involves the use of embedded AI chips that can perform complex computations, such as pattern recognition, decision-making, and machine learning tasks, without relying on cloud connectivity. This shift not only enhances real-time decision-making but also addresses pressing concerns around data privacy and security.

The role of edge AI in creating sustainable, energy-efficient systems that can operate autonomously in diverse environments, from rural communities to densely populated urban centres (that directly address global challenges), was explored at a workshop at the Summit.

Presentations can be found at the [Edge AI Workshop](#) website.

Some of the key advantages of Edge AI include:

1. **Reduced latency**

Edge AI minimizes response times by processing data locally rather than relying on cloud servers. This is crucial for applications like autonomous vehicles.

2. **Bandwidth conservation**

Autonomous vehicles are essentially moving data centres on wheels, generating massive amounts of information every second. A single self-driving car can produce terabytes of data daily from its cameras, LiDAR, radar, GPS units, and IoT sensors. This raw information needs to be processed immediately to ensure the car can navigate safely and effectively. For example, autonomous vehicles may need to process between 3 Gbit/s to 40 Gbit/s of sensor data, depending on their level of autonomy. Sending all of this raw data to the cloud would not be practical and create unnecessary costs. Edge AI allows cars to filter, and process data locally, ensuring that only the most valuable insights are transmitted to the cloud. This makes fleet management more efficient, as cloud systems receive processed summaries rather than endless streams of raw sensor data.

3. Offline capability

Edge AI enables robots to function autonomously in environments with limited or no internet connectivity. This is essential for applications like disaster response or space exploration. This enhanced reliability is particularly crucial for mission-critical applications where continuous operation is necessary, even in remote or disconnected environments. Edge AI ensures high availability for devices by enabling them to operate autonomously without relying on continuous internet connectivity or cloud-based services.

4. Enhanced data security and privacy

Processing sensitive information locally on devices using edge AI enhances data security by reducing the risk of exposure or attacks during transmission to cloud servers. This approach keeps critical data within the device, minimizing the chances of unauthorized access, data breaches, or interception.

5. Energy efficiency

Recent research shows that shifting neural processing from CPUs and GPUs to specialized hardware AI processors can greatly reduce power use in edge devices. By performing computations locally and sending less data over the network, edge AI improves power efficiency and reduces dependence on energy-hungry cloud servers. This lowers the carbon footprint of edge devices by minimizing the I/O operations needed for cloud AI applications.

6. Real-time performance

Edge AI delivers high-performance computing on local devices, instantly processing data and running machine learning and deep learning algorithms. Unlike cloud processing, it works in milliseconds, making it perfect for real-time applications like defect detection in production lines and abnormal behaviour detection in security systems.

13.2 Where is edge AI needed?

Edge AI is particularly important in industries where timing, precision, and adaptability are crucial. In manufacturing, traditional systems rely on static programming and centralized control, while edge AI allows for more dynamic adaptation. For example, the automotive

industry is using edge AI in processes like closed-loop gluing. Here, instead of sending data to a distant server for analysis, the system adjusts in real-time, optimizing the process as it happens.

In healthcare, AI can reduce workload for clinicians and lead to new innovative solutions, for instance with the use of voice-enabled assistants to help keep sensitive data local. It is possible to compress AI models and fine-tune them for healthcare applications at the edge. Moreover, in healthcare, Edge AI is being used to enhance diagnostic tools. Medical devices equipped with AI that can process data locally, such as ultrasound machines or portable scanners, enable doctors to receive real-time feedback, which can be critical in clinical settings. Instead of waiting for a centralized system to provide analysis, doctors can make quicker, more informed decisions, potentially improving patient outcomes.

Robots that process data locally are not only faster but also more responsive to immediate environmental changes. This can lead to more fluid, human-like interactions and decision-making, without the delays associated with cloud-based processing.

Yet, while it brings many advantages, adopting edge AI requires careful consideration of the specific needs and constraints of each use case.

The current generation of embedded and edge systems serves industries such as aerospace, defence, industrial, medical, automotive, and telecommunications. Much of these systems' application logic relies on statically compiled code or dynamically loaded libraries, with minimal embedded intelligence.

13.3 Edge AI hardware

With the integration of diverse sensors into devices and rapid advancements in nanoscale chip technology, significantly more data can be captured and processed directly at the edge, without needing to route to the cloud. Additionally, silicon manufacturers are embedding AI capabilities within systems-on-chip (SoCs), which allows compact AI runtimes and enables frameworks to run on a variety of processors, including NPUs, GPUs, FPGAs, and ASICs. These capabilities span silicon architectures such as x86, Arm, and RISC-V. They require real-time operating systems and Linux, which are predominant in embedded and edge systems, to support AI at the edge.

It is expected that the intelligent edge will support diverse silicon and the applications that run on it. That means new generations of distributed computing environments must be optimized for power and compute efficiency, compact AI runtimes and frameworks, hardware virtualization, and AI Ops for AI-enabled applications.

AI models capable of running efficiently on edge devices need to be reduced considerably in size and compute while maintaining similar reliable results. This process, often referred to as model compression, involves advanced algorithms like neural architecture search (NAS), transfer learning, pruning, and quantization. Model optimization should begin by selecting or designing a model architecture specifically suited to the device's hardware capabilities, then refining it to run efficiently on specific edge devices. NAS techniques use search algorithms to explore many possible AI models and find the one best suited for a particular task on the edge device. Transfer learning techniques train a much smaller model (the student) using a larger model (the teacher) that's already trained. Pruning involves eliminating redundant parameters

that don't significantly impact accuracy, and quantization converts the models to use lower-precision arithmetic to save on computation and memory usage.

Some of the most popular AI models for vision applications—are designed to be extremely efficient at these calculations. But in practice, these models do not always run well on the AI chips inside our phones or smartwatches. This is because real-world performance depends on more than just math speed—it also relies on how quickly data can move around inside the device. If a model constantly needs to fetch data from memory, it can slow everything down, no matter how fast the calculations are.

13.4 Edge AI use case implementations

13.4.1 Addressing untreated hearing loss

AI is rapidly transforming hearing technology. femtoAI addressed some of the biggest factors contributing to the prevalence of untreated hearing loss with SPU-001 processors and deep learning noise reduction algorithms by a.) solving the number 1 reported performance issue with hearing aids, speech quality in noise, b.) reducing the stigma by enabling smaller, more discrete form factors, and c.) reducing cost by bringing technologies to low cost OTC hearing aids for an affordable price.

femtoAI developed an AI-driven chip technology for hearing aids, earbuds, and other audio devices, using a method that mimics the brain's efficiency to enable AI processing on low-power devices. The AI chip optimizes speech separation and noise reduction while maintaining battery efficiency – critical for hearing devices that operate within strict power constraints. Hearing aids are becoming an arena for cutting-edge AI innovation. With AI-driven real-time speech enhancement, modern hearing aids can now improve the clarity of conversations in noisy environments, addressing one of the biggest barriers to adoption. As AI capabilities advance, hearing aids are shifting from simple preset sound adjustments to adaptive, real-time sound processing, providing a more personalized and effective hearing experience.

13.4.2 MountAIn Edge AI infrastructure in remote environments

In recent years, developments in IoT, AI, and satellite communications have been providing data intelligence in ways that previously seemed inconceivable. But these technologies remain somewhat siloed. While the potential opportunities in bringing these technologies together are significant, challenges related to connectivity, power availability, and compact, embedded design remain.

MountAIn has been able to address these challenges in their AI end-to-end infrastructure. At its core is the IBEX device, a "digital watchman" specifically trained to monitor unique use case parameters. As an advanced IoT imaging device, it has embedded AI processing while also acting as a LoRa® gateway for additional data inputs. IBEX operates without a power source or Wi-Fi, relying instead on solar energy. When required, events-based short messaging alerts are sent to a user via satellite connectivity. Its versatility makes it a promising solution for a wide range of scenarios, from managing urban assets to monitoring remote natural environments. It is being used to automate monitoring the risks of fire hazards in forests and the olive fruit fly in precision agriculture.

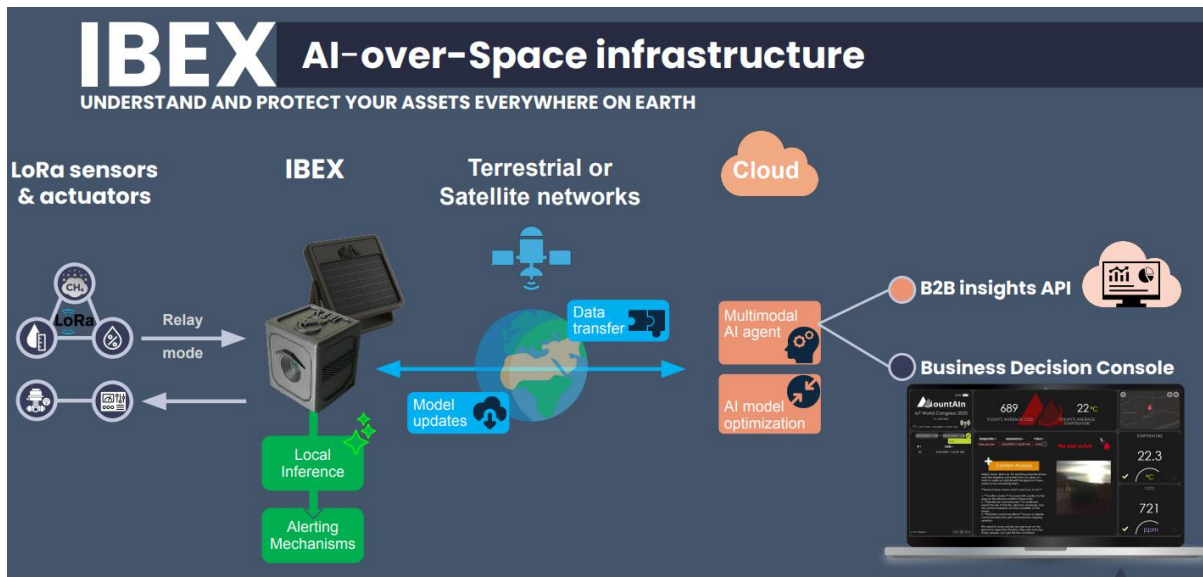


Figure 42: MountAin IBEX

13.4.3 Building trustworthy local AI for healthcare at the edge

Embedl shared insights about how AI can be both powerful and private by operating securely on local servers and edge devices rather than relying solely on the cloud by giving the example of voice-enabled assistants in the healthcare domain.

13.4.4 Implementing edge AI in mission-critical industries

One important example of the need for edge AI is in the manufacturing sector. As intense competition, skilled labour shortages, inflation, and stringent quality expectations rise, manufacturers are increasingly turning to automation and more collaborative operation between humans and robots.

This is leading to new health and safety concerns, such as the risk of physical collision, the increased need for monitoring and oversight of often repetitive tasks, and the risk of system failures or unexpected behaviour.

Wind River explored how organizations can balance innovation with responsibility – ensuring that AI deployed at the edge not only advances business and operational goals but does so in a way that is safe and accountable. The Wind River solution is designed to prevent injuries by minimizing human error and enhancing workplace safety in environments with increased automation. Wind River’s solution comprises a solution for real-time, context-aware edge AI which is pre-integrated with market-leading hardware to simplify the use of AI to improve industrial operations.

This approach powers critical applications – including computer vision, sensor analytics, industrial automation, and security – across a wide range of industries, such as manufacturing, healthcare, logistics, and energy.

A camera monitors the manufacturing floor and the AI system analyzes the image to ensure workers are equipped with the correct safety gear, such as hard hats and vests. When workers have the proper protective equipment, the system enables robots in the area to become operable.

If someone removes a hard hat, for example, the system detects the change and instantly stops robots from operating.

13.4.5 From classroom to community: Five years of TinyML academic network impact

Lessons learned from building and scaling the network of tinyML community highlight how collaborative efforts among educators, students, and industry partners have fostered innovation, inclusivity, and real-world impact. Challenges faced, such as resource limitations and diverse educational contexts, were addressed through partnerships and open-source approaches.

13.4.6 AI inference chip

AI development is typically divided into two stages: training, which demands massive datasets and compute, and inference, where trained models are deployed to solve real-world problems. As AI adoption broadens, cloud-based inference is quickly taking centre stage. According to the International Data Corporation (IDC), cloud-based inference accounted for 58.5% of AI computing power in 2022 and is projected to hit 62.2% by 2026. It is forecasted that AI inference compute demand will grow over 80% annually, potentially surpassing training as the primary driver for data centre expansion.

Intellifusion developed AICB (AI Computation Block), an innovative architecture for edge AI which is enabling AI inference chips to meet the computational demands of various devices. The diagram below shows the evolution of Intellifusion's inference chips.

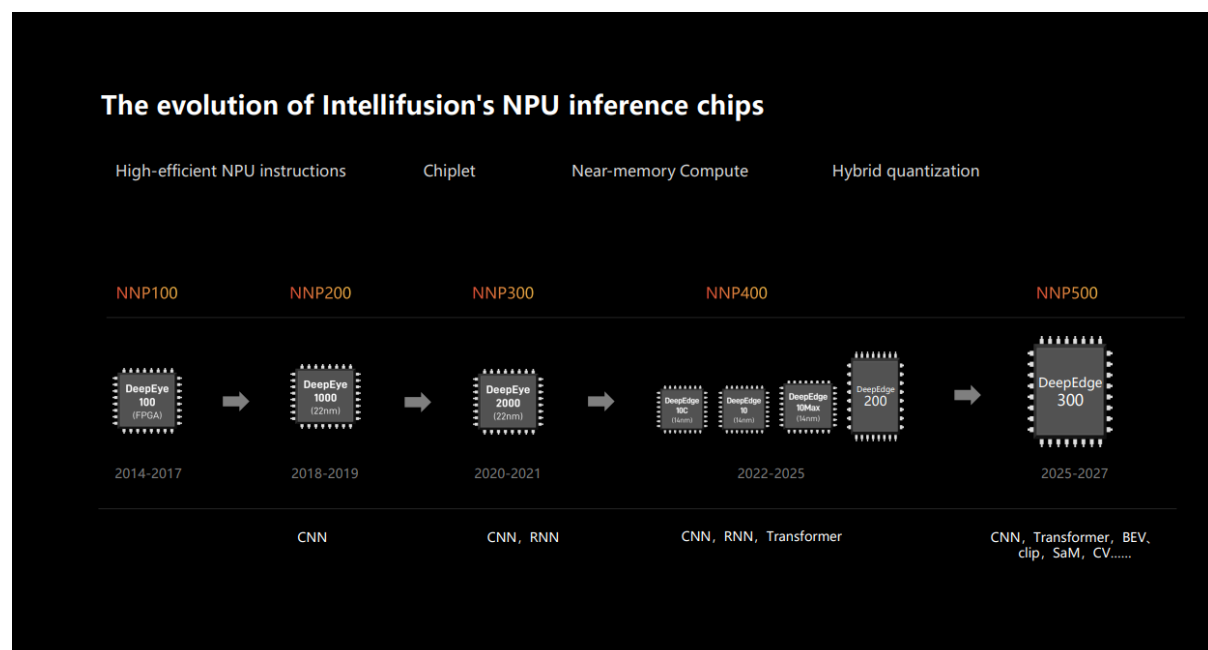


Figure 43: Evolution of NPU reference chips of IntelliFusion

13.4.7 Day 0 support for novel models enabling instantaneous deployment of edge AI

AI compiler technology has become mission-critical for deploying AI models at scale on edge devices. The usual trend is for solutions to build on a legacy software stack that interprets AI models but does not allow for compilation. This limits the application of the technology for AI models, such as language models, as they are not compatible with the existing technology stack.

AI compilation optimizes model execution at different levels of abstraction – known as "intermediate representations" - that represent specific features of the execution of an AI workload. The compiler translates the AI model through various intermediate representations to a low level that is close to the target hardware. Using AI compilation instead of model interpretation allows users to adapt to the constant stream of new AI models.

Roofline's flexible SDK enables models to be imported from any AI framework such as TensorFlow, PyTorch, or ONNX. It enables deployment across diverse hardware, covering CPUs, MPUs, MCUs, GPUs, and dedicated AI hardware accelerators - all with just a single line of Python. Roofline's flexible SDK that enables day-0 support for PyTorch models found on Hugging Face, which solves the challenge of the inability to support emerging new AI models and layer types, due to specific software features not compatible across systems.

13.4.8 Defining the boundaries of AI from fundamental limitations to resource constraints

CSEM mentioned the challenges and solutions in deploying AI at the edge, focusing on the limitations posed by neural scaling laws and the memory wall. The need for hardware-aware AI design, efficient feature selection, and adversarial training to overcome resource constraints while maintaining performance was highlighted. The proposed strategies enable scalable, private, and energy-efficient AI solutions tailored for edge computing environments.

13.5 Key takeaways

1. TinyML and edge AI feature low energy cost, comparatively cheap price, and high functionality because of being real-time, task distribution, and cloud-free characteristics.
2. TinyML and edge AI are growing in terms of scale and domain, despite the challenges.
3. Regulation and policy support for the larger-scale use of TinyML are needed. Capacity building, infrastructure, TinyML and edge AI embedded strategies, pilot projects, and regional collaboration are key components to improve the adoption and deployment of TinyML and edge AI.

14 Navigating the intersect of AI, environment and energy for a sustainable future

In an era in AI is rapidly reshaping industries, understanding its environmental impact and energy dynamics becomes paramount for steering towards a sustainable future. This multidisciplinary workshop aimed to unravel the complex relationship between AI, the environment, and energy consumption; spotlight innovations driving AI environmental efficiency; explore AI's transformative potential for environmental efficiency in several sectors; and deliberate on the pivotal role of standards, policies, and regulations.

This event aligned with ITU's Green Digital Action initiative, reinforcing ITU's commitment to promoting digital innovation, standardization, and global collaboration to foster sustainable AI development while ensuring the ICT sector minimizes its environmental impact and maximizes its transformative potential.

Presentations for this workshop can be accessed [here](#).

14.1 Understanding AI's environmental impact

This session provided a comprehensive examination of the environmental footprint of AI systems, emphasizing the resource intensiveness of the AI system lifecycle from data collection and model training to deployment and inference.

Research from Harvard University on the environmental impact of hyperscale data centres in the US found that 403 such data centres are responsible for more than 52 million metric tons of CO₂ annually, representing 1.1 per cent of US electricity-related emissions in 2023. The carbon intensity of hyperscale data centres is 52% higher than the US average, highlighting the need for targeted, location-specific strategies. The disproportionate effects on marginalized communities were also highlighted, underlining the need for data-driven planning tools that integrate public health and equity concerns. A call was made for a holistic approach that goes beyond CO₂ emissions to also account for other factors such as air quality.

Nvidia highlighted significant energy efficiency gains across its AI stack. For example, LLM inference energy efficiency has improved 100,000 fold over the past decade, while liquid cooling systems have reduced water usage by a factor of 300. AI contributes to electricity load growth but still represents a small share (~0.3%) of global electricity use and is increasingly powered by clean energy, according to IEA World Energy Outlook 2024 and IEA Energy & AI Report 2025.

The importance of standards and sectoral cooperation to mitigate AI's environmental and material impacts through the standardization work of ITU-T Study Group 5 (Environment,

electromagnetic fields, climate action, and circular economy) were shared. Work addressing the need for a stable and standardized methodology for measurement include:

- Draft Recommendation ITU-T L.1472 on GHG emissions and energy consumption of the ICT sector database and pilot project on the implementation of the standard
- Draft ITU-T L.EnvAI guidelines on AI systems environmental impact
- Guidelines for cities to achieve carbon Net Zero through digital transformation

The ongoing Swiss-based lifecycle assessments of Mistral AI models, evaluating emissions from both training and inference over a two-year period, will also help inform standardized environmental impact assessments across countries. Recommendation ITU-T L.1450 provides "Methodologies for the assessment of the environmental impact of the information and communication technology sector" and work is underway on a new ITU standard to assess the environmental impact of artificial intelligence systems.

France's national strategy for sustainable AI, focusing on multistakeholder collaboration among government, academia and industry, is developing actionable initiatives with international partners, such as the recently launched Coalition for Sustainable AI.

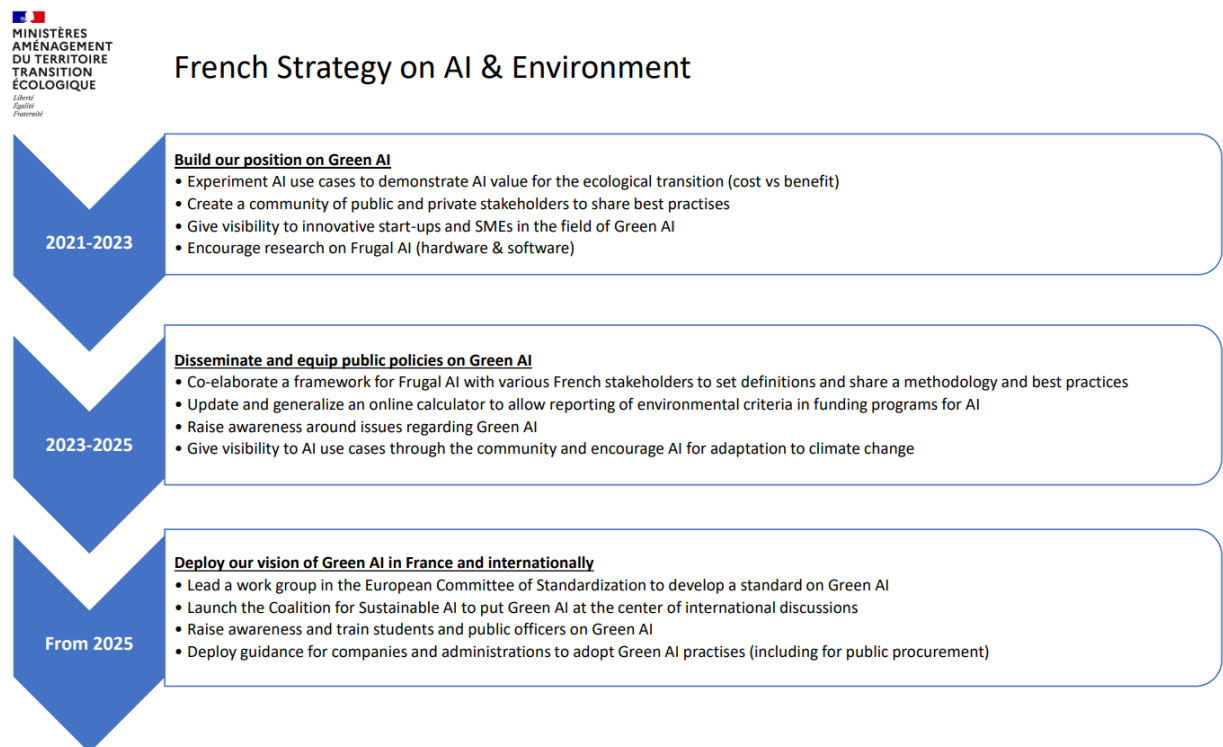


Figure 44: France strategy on AI and environment

HSBC focused on challenges in current AI footprint assessments, noting issues such as:

- The lack of standardized tools and lifecycle boundaries, which hinders comparability
- Underreporting of Scope 3 emissions (e.g., hardware, e-waste)
- Omission of water usage and cooling overheads

- Over-reliance on estimates rather than observed telemetry data

The ITU report on [Measuring What Matters: How to Assess AI's Environmental Impact](#) was also launched at the event. This report offers a comprehensive overview of current approaches to evaluating the environmental impacts of AI systems.



Figure 45: ITU report on measuring AI's environmental impact

Key takeaways

- AI's environmental impact is real, complex and growing, not only in terms of energy and emissions, but also with regard to water usage, material intensity and equity considerations.
- Significant progress has been made in energy efficiency, but further efforts are needed to standardize measurements and expand impact assessments to include the full lifecycle and broader ecological factors.
- There is a clear need for harmonized global standards and multistakeholder cooperation, especially involving regulators, private sector, and UN agencies, to guide AI deployment in alignment with climate goals.
- Opportunities exist for ITU-T Study Group 5 to advance work on AI-specific standards, particularly in:
 - Model-level environmental accounting
 - Data centre efficiency benchmarking
 - Lifecycle reporting frameworks
 - Water use and biodiversity metrics
- The role of regulation was debated, with some speakers advocating for more stringent regulation, while others questioned its effectiveness, arguing that it may not provide sufficient incentives.

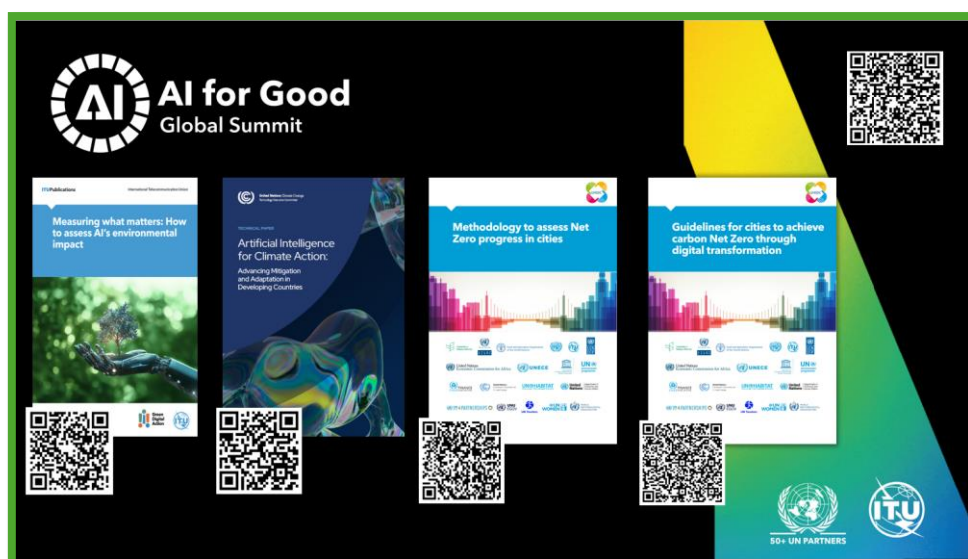


Figure 46: Reports announced at the event

14.2 Innovations in environmentally efficient AI

This session explored technological innovations aimed at reducing AI's environmental footprint across the hardware-software stack, sharing perspectives on energy-efficient AI models, neuromorphic computing, sustainable data centre operations, and the role of green energy.

University College London highlighted the daily impact of energy savings and the importance of optimizing large models, while using small models when appropriate. The importance of designing for efficiency from the outset was emphasized, with scenarios showing that task-specific small models can achieve over 90% energy savings compared to large, multipurpose models. Emerging architectures such as Mixture of Experts, retrieval-augmented generation, neurosymbolic AI, and brain-inspired designs were discussed as promising pathways toward sustainability.

Current digital computing architectures pose theoretical and practical limitations for sustainability, leading participants to consider that spiking neural networks and neuromorphic computing could be biologically inspired alternatives that drastically reduce energy use while enhancing reliability.

The "Green AI" movement also emphasizes computational efficiency and transparency. While large models like PaLM require massive amounts of computing resources, most environmental impact comes from inference, not from training. The Hebrew University of Jerusalem noted that inference operations account for 80-90% of all AI computation and are run billions of times per day. There is a need for the community to report compute budgets and match model complexity to task difficulty, considering that LLMs are not the solution for every problem.

Google's end-to-end sustainability strategy includes:

- Model optimization (e.g. quantization, pruning, and knowledge distillation)
- Custom hardware (e.g. Ironwood TPU, 30x more efficient than its 2018 predecessor)
- Energy-efficient data centres (6 times more compute power than 5 years ago)
- A target for 24/7 carbon-free energy by 2030, through partnerships with geothermal and small modular reactor (SMR) developers
- Water replenishment goals and AI-driven cooling optimization (e.g. DeepMind's 40% reduction in cooling energy)

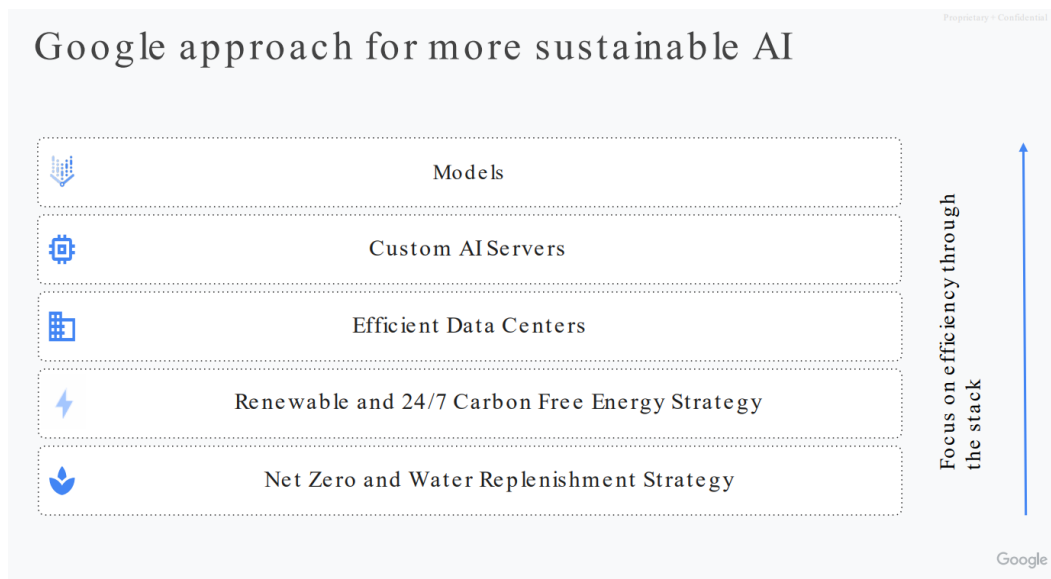


Figure 47: Google approach for sustainable AI

Some of the key takeaways are summarized below:

- a) The energy efficiency of AI is a challenging problem to solve.
- b) Energy and inference matter: Alongside the energy demands of large-scale model training, attention should also be paid to sustainable inference – the true computational bulk of AI.
- c) Hardware-software co-design is critical: Emerging architectures (such as neuromorphic, brain-inspired, or task-adaptive) can help cut energy consumption dramatically, as shown in the use of small models and analogue systems.
- d) Green AI should become the norm, not the exception – through better compute reporting, evaluation metrics, and research funding incentives.
- e) Standards and benchmarks are needed for:
 - Model-level energy reporting
 - Inference efficiency metrics
 - Lifecycle footprint of AI chips and cooling systems
 - Frugal AI design methodologies

f) Opportunities for ITU:

- Support the development of standardized AI energy efficiency frameworks, potentially building on the ongoing work of ITU-T Study Group 5 and consider developing KPIs to score the environmental efficiency of AI systems..
- Consider developing new technical specifications on sustainable AI deployment, data centre optimization, and model auditing.



Figure 48: UNESCO Report on Smarter, Smaller, Stronger: Resource Efficient Generative AI and the future of digital transformation

14.3 AI-driven environmental sustainability across industries

The intersection of AI and environmental sustainability across diverse sectors was highlighted through examples of practical implementations and theoretical frameworks aimed at aligning AI development with climate and ecological goals.

Examples included:

- CodeCarbon is an open-source tool that tracks CO₂ emissions from code execution and provides actionable insights for researchers, developers, and data scientists. AI's footprint extends beyond CO₂ emissions to water usage and hardware manufacturing.

Scope 1–3 impacts include:

- Scope 1: direct emissions from cooling data centres
- Scope 2: emissions from electricity used
- Scope 3: chip production and infrastructure

Concrete mitigation techniques such as model pruning, quantization, distillation, and optimized training locations, while encouraging users to "use the smallest model for your need" and extend the lifespan of hardware, could help address these challenges.

- Discussions on the sustainability risks of generative AI noted that hardware growth cannot match model scale, leading to the need to consider:

- Cost-effective, generalizable models
- Repurposing pretrained models for new use cases
- Open-source ecosystems to ensure transparency and democratized development
- Greater efficiency through edge computing and modular designs
- ZTE emphasized that global collaboration and human-centred principles are necessary to balance innovation with energy constraints. ZTE's vision of a three-tier strategy for green AI includes:
 - Efficient infrastructure: Enhancing computing and cooling through distributed architecture and energy-aware networking, in particular when developing high-speed networks.
 - Intelligent empowerment: Optimizing AI through multimodal learning and model compression techniques.
 - Efficient implementation: Deploying AI across sectors such as smart grid, disaster warning, green finance, and personalized energy tracking.
- PrevisIA uses AI to forecast deforestation in the Brazilian Amazon by analysing satellite imagery, historical deforestation patterns, topography, water bodies, and socioeconomic data, with a focus on detecting unofficial roads – a key predictor of forest loss. The initiative integrates three pillars: AI-driven road detection via Sentinel-2 imagery, predictive modelling with geospatial dashboards, and collaborative enforcement with state prosecutors. Partnering with Public Prosecutor's Offices in Pará, Amazonas, Acre, and Amapá, PrevisIA achieved 73% forecast accuracy (2021–2024) within a 4km radius. Its reports enable embargoes, fines, and legal action against illegal deforestation, demonstrating AI's potential for proactive environmental protection. The system integrates remote sensing and geospatial AI to assess risk factors like proximity to roads, showing that 95% of deforestation occurs within 5.5 km of roads. PrevisIA supports enforcement actions (e.g. fines and licensing suspensions) and empowers authorities with real-time alerts.
- UNFCCC launched the TEC's technical paper on AI for Climate Action, which focuses on applications in developing countries, including:
 - Mitigation with energy optimization, emissions tracking, and renewable energy.
 - Adaptation with early warning systems and localized services for smallholder farmers.
 - Policy with practices to accelerate and transfer climate actions to developing countries.
 - Governance with responsible AI frameworks to mitigate risks in fragile contexts.

Some of the key takeaways are highlighted below:

- a) AI's environmental impact is multidimensional, encompassing emissions, water, material extraction, and infrastructure. Inference and deployment now dominate energy use, not just model training.
- b) Tools (including open-source tools) are being developed to offer practical pathways for emissions estimation and reduction and can help align developers' daily work with sustainability goals.
- c) Hardware limits are a real constraint on AI scalability, calling for innovations for greater model efficiency in the interest of sustainability.
- d) Sectoral applications (e.g. deforestation prevention, smart energy, and public health) showcase AI's potential for impact — if implemented responsibly.
- e) Shared global progress is essential to environmental sustainability, highlighting the importance of capacity-building in developing countries.
- f) Opportunities for ITU-T Study Group 5;
 - Support the development of standardized metrics and tools (e.g. model-level energy benchmarks, and lifecycle impact metrics).
 - Collaborate on guidelines for responsible AI deployment in climate-sensitive and resource-constrained environments.
 - Facilitate cross-sector knowledge-sharing and open innovation platforms to promote scalable, low-impact AI solutions.

14.4 Standards, policies and regulations for sustainable AI and environment

Regulatory frameworks and international standards can help guide the sustainable deployment of AI in the energy and ICT sectors. It will also be important to consider data transparency, institutional cooperation, and the alignment of national policies with global climate goals, as well as clearly define the various stakeholders' roles in addressing the environmental impacts of AI.

Colombia's data centre regulation project aims to improve the transparency and environmental accountability of digital infrastructure as part of the country's effort to collect data on energy consumption and GHG emissions. Recognizing the importance of political will private-sector participation, a call was made for regional cooperation and harmonized standards, particularly for reporting frameworks that could be adopted across Latin America.

Brazil offered a national perspective in the lead-up to COP30, which it will host. Brazil is committed to environmental protection and climate diplomacy, given the centrality of sustainable digital transformation and AI governance in shaping Brazil's COP30 priorities. From Brazil's perspective, there is a need to consider embedding AI energy regulation into the broader UNFCCC agenda and multilateral institutions like ITU could help to scale support for capacity-building in emerging economies.

The African Telecommunication Union (ATU) underlined the regulatory and infrastructural gaps in Africa that challenge sustainable AI deployment and called for fit-for-purpose standards that reflect the realities of developing and least developed countries, especially where digital access is still limited. ATU noted that the ICT sector can play a dual role in accelerating access to energy through smart grids and reducing its own footprint, but only if regulators are equipped with the right tools and data.

Considering the work of France and the EU on digital environmental regulation, referencing the European Data Act and energy performance monitoring of data centres, binding obligations for AI developers and data centre operators, including disclosure of energy intensity, carbon footprint, and water usage were deemed to be important. The discussion also noted the importance of alignment between national regulators and global standard-setting bodies like ITU, recalling the importance of data measurement at the international level.

Some key takeaways include:

- a) Data access remains a bottleneck. Several panellists echoed the need for greater transparency on energy consumption and GHG emissions, with this data being essential for developing effective regulations and standards.
- b) Considering the dual nature of AI as a tool for energy optimization as well as a source of energy demand, regulators could consider systems-level regulatory approaches that integrate AI policy with climate and digital infrastructure policy.
- c) Developing countries need support, and capacity building for policymakers and regulators will be important, particularly in Africa and Latin America.
- d) Calls were made to accelerate the adoption of relevant standards from standards bodies such as ITU.
- e) The session framed COP30 as a strategic moment to embed digital sustainability and AI governance into global climate action.
- f) **Opportunities for ITU-T Study Group 5**
 - Facilitate the creation of a global reporting framework for AI-related energy and emissions data, in cooperation with national regulators.
 - Support governments in adapting ITU standards to local contexts, with technical assistance, training, and multistakeholder engagement.
 - Collaborate with regulators to co-develop policy toolkits that guide sustainable digital infrastructure development and AI deployment.

15 Future Networked Car Symposium 2025

The Opening ceremony set the tone for the day by emphasizing the strategic role of AI in shaping the future of safe, sustainable, and inclusive mobility. Speakers highlighted the vital need for collaboration among UN bodies, standards bodies, industry leaders, and the research community to address the complex challenges of AI adoption in the automotive sector. Speakers also underscored the importance of standardization, data sharing, and road safety to sustainable development. The opening session launched the symposium with a strong call for coordinated action and shared global responsibility. All the presentations can be accessed [here](#).

15.1 How AI can improve mobility for a better future

AI is revolutionizing the mobility landscape, potentially offering unprecedented benefits in safety, efficiency, and user experience. A primary motivation for autonomous vehicles (AVs) integrating AI and robotics is improving road safety, as human error is the cause of the majority of traffic crashes. This session discussed how AI can help optimize traffic management, enhance vehicle safety through predictive maintenance and collision avoidance systems, and personalize mobility services to meet individual needs.

15.2 The revolution of AI automotive systems infrastructure

AI is touching all points in the automotive industry from vehicle design, manufacturing, marketing, and hiring to the functioning of in-vehicle systems. In-vehicle systems, too, are being transformed from digital assistants and contextually enhanced experiences to automated and semi-automated driving and collision-avoidance technology.

This session focused on the technical infrastructure required to deliver AI benefits in automotive systems and on the development of guidelines, standards, and regulations for automotive AI.

For example, BMW's Panoramic Vision is a new heads-up display that projects information, such as vehicle speed, onto windshields. It uses AI to track the driver's eye movements to ensure that information always appears in a clear, easy-to-see location that will not interfere with the view of the road. The most important information appears in a darkened section at the base of the windshield, while less important data are projected slightly higher onto a clear section. The technology can automatically determine when certain information should be displayed in a priority spot. For instance, it prominently displays navigational information to help with parking when a driver is looking for or entering a parking spot.

15.3 AI automotive applications ethics and their human elements

As AI becomes integral to automotive systems, ethical considerations surrounding human oversight, data privacy, and equitable access come to the forefront. This session explored how AI can support human drivers without overwhelming them, the importance of prioritizing certain types of vehicle data for sharing, and the need for ethical fail-safes in case of system malfunctions.

15.4 MoU between ITU and SAE

During the symposium, ITU and SAE International signed a Memorandum of Understanding to establish a framework for cooperation in areas of shared interest related to intelligent transport systems and mobility technologies. The collaboration aims to promote standardization for vehicular communications, enhance stakeholder engagement, and contribute to digital transformation and the UN Sustainable Development Goals.



Figure 49: Left: Seizo Onoe, Director of the Telecommunication Standardization Bureau (TSB), International Telecommunication Union (ITU). Right: Christian Thiele, Senior Director, Global Ground Vehicle Standards, SAE International

Standardized frameworks and interoperable infrastructure play a critical role in advancing safe, ethical, and inclusive AI-driven transportation systems. The dialogue during the symposium reinforced the need for continuous alignment between technical innovation and regulatory oversight, with human-centric values at the core. Looking ahead, stakeholders recognized the importance of ongoing engagement between industry, government, and research communities to ensure that AI in mobility delivers sustainable, equitable, and trustworthy solutions.

15.5 The role of standards in shaping autonomous mobility

While the symposium covered a wide range of topics related to AI in mobility, a recurring theme throughout the discussions was the critical importance of standards. As autonomous driving technologies advance, harmonized standards will be essential to safety, interoperability, and public trust.

Various ITU-T study groups contribute to intelligent transport systems (ITS):

- [ITU-T Study Group 2 \(Operational aspects\)](#) works on numbering and addressing for in-car emergency communications, including the assignment of international numbering ranges for services like automatic emergency calls from vehicles.
- [ITU-T Study Group 12 \(Performance, QoS and QoE\)](#) works on speech and audio quality in vehicles, including hands-free communication and automatic speech recognition, with the aim of enhancing driver safety and minimizing distractions.

- [ITU-T Study Group 17 \(Security\)](#) works on security aspects of ITS including secure over-the-air software updates, intrusion detection systems, and guidelines for vehicle-to-everything (V2X) communication and electric vehicle charging.
- [ITU-T Study Group 20 \(IoT, digital twins and smart cities\)](#) works on digital twins for traffic systems, roadside sensing for autonomous driving, and emergency response systems.
- [ITU-T Study Group 21 \(Multimedia, content delivery and cable TV\)](#) works on vehicular multimedia and AI for automated driving, including standards for vehicle gateways, infotainment systems, and V2X communication.

In addition, all standards bodies working on ITS are represented in the ITU-led [Collaboration on ITS Communication Standards \(CITS\)](#). The CITS Expert Group on Communications Technology for Automated Driving is developing requirements for next-generation vehicular communications. Its working groups are addressing critical challenges such as:

- Automated merging into congested lanes ([WG1](#)), aiming to define communication protocols and safety requirements for vehicles with active automated driving systems. [\[CITS EG-ComAD/WG1\]](#)
- Advanced emergency braking systems ([WG2](#)), including protection for vulnerable road users, in alignment with emerging global safety regulations.

The Future Network Car Symposium 2025 reaffirmed the need for ongoing collaboration among standards bodies, industry players, and policymakers to build a future of mobility that is safe, inclusive, and sustainable.

16 AI readiness workshop

16.1 Introduction

The AI Readiness workshop during the AI for Good Summit in May 2024 served as a foundational platform for discussion on these questions, where the [report](#) titled "Preliminary Analysis Towards a Standardized Readiness Framework" was published. This preliminary report examined multiple AI use cases and identified six key readiness factors:

- **Data:** Accessibility and quality of datasets for analysis of AI applications.
- **Research:** Collaboration between domain-specific and AI research communities.
- **Deployment support:** Infrastructure and ecosystem readiness for AI deployment.
- **Standards:** Ensuring trust, interoperability, and compliance.
- **Open source and code:** Enabling rapid adoption through an open developer ecosystem.
- **Sandbox environments:** Platforms for AI experimentation and validation.

To advance these discussions, version 1.0 of the ITU AI Readiness [report](#), along with a Call for Engagement towards the ITU AI Readiness Plugfest, was launched by ITU and the Kingdom of Saudi Arabia during the [2024 GAIN Summit](#).

The ITU AI Readiness Plugfest is an initiative to explain the AI readiness factors to various stakeholders and provide a platform for stakeholders to "plug in" their regional AI readiness factors, such as data accessibility, AI models, compliance with standards, toolsets, and training programs. Additionally, a Technical Advisory Committee, composed of experts invited through AI for Good Impact initiative, provides strategic guidance and feedback on AI readiness projects. To study the sandbox environments and influence AI readiness, cloud credit support is provided to selected projects, further facilitating the development and deployment of AI solutions in real-world applications.

The AI Readiness Workshop at the AI for Good Global Summit 2025 complemented and enhanced the AI readiness studies with a strategic set of assessment indicators which will assist stakeholders, especially in developing countries, in evaluating and improving their AI readiness status. The workshop also discussed case studies, priority areas for attention and resource investment and improving global AI capacity building and fostering opportunities for international collaboration based on existing assessment mechanisms.

By fostering collaboration among global stakeholders, industry leaders, and researchers, this workshop aimed to support the standardization of AI readiness evaluation and accelerate AI adoption across diverse domains. The workshop also identified the next steps leading to the release of an updated AI Readiness Report, incorporating new findings, industry feedback, and key developments from the Pilot AI Readiness Plugfest.

Presentations can be found at [AI Readiness Workshop](#).

16.2 Pilot AI readiness plugfest presentation and partner sharing

Some of the main projects that were presented are shown below:

- i. Saudi Data and AI Authority presented the strategy of the data and information building of Saudi Arabia, with a use case of EYENAI for diabetic retinopathy. EYENAI is an AI-powered medical solution for accurate diabetic retinopathy detection and diagnosis. It utilizes advanced analytics and intelligent algorithms to streamline the examination process, addressing the challenges posed by costly, limited resources, and time-consuming traditional methods.
- ii. CAICT highlighted AI model adoption choices and preferences in various domains, and the potential actions for the next version of the AI Readiness Report and Framework, such as creating assessment tools, developing measurable criteria, and metrics.
- iii. Global Green Growth Institute shared the framework and model its built to analyze the co-benefits of the UN Sustainable Development Goals and provide scenario time-series prediction visualization.
- iv. Lenovo presented its ITU-supported project of AI-Powered Human-like Avatar for Sign language in Brazil and its impacts.
- v. Umgrauemeio (1.5 degree) presented its project focusing on wildland fire management.
- vi. Pontificia Universidad Católica de Chile presented the progress on its Latin America AI Index (ILIA), including the large amount of data it has collected for assessing and ranking. Inputs from ILIA will support the development of the ITU AI Readiness toolkit.
- vii. MGIMO AI Center presented AI readiness in international trade from a process-centric point of view. Their input to the ITU AI Readiness Framework considers the importance

of finding the right place to apply AI by inputting domain-specific workflows identified by authoritative standards in the relevant field, and quality-related risks relevant to trading AI goods and services.

- viii. DevelopMetrics emphasized the importance of preserving local data and forming a platform with a standardized data format for future reference.
- ix. The University of Applied Sciences and Arts of Southern Switzerland presented its project on drone-supported wheelchairs with a special focus on lawful, ethical, and robust use of AI.
- x. Australia's University of Victoria shared its project using AI to provide companions for the visually impaired population using LLMs and vision-language models. It was observed that the evaluation of an AI application is key to assessing AI readiness.

16.3 Outcomes

A key outcome was planning the next steps toward a revised AI Readiness Report, which will incorporate new methodologies, feedback from industry, and complement the lessons learned from the pilot AI Readiness Plugfest with the assessment indicators. Building on these insights, ITU will take AI readiness discussions forward by:

- Enhancing AI readiness evaluation methodologies by studying the characteristics of AI integration in various domains.
- Exploring pathways for stakeholders to contribute to a standardized AI readiness framework.
- Engaging the Technical Advisory Committee to provide strategic guidance and feedback.

The ITU AI Readiness Framework aims to enable countries to conduct the AI readiness self-assessment. ITU is calling for engagement from experts in different domains to design and refine the key factors and relevant indicators to deliver a toolkit. Learnings from plugfest project owners are essential as they bring in regional and domain-specific points of view. Plugfest presentations are supporting the identification of common metrics for measuring AI readiness.

17 AI for agriculture – shaping standards for smart food systems

AI has the potential to help transform agriculture. By fostering trust and sound governance, AI-driven food systems can become smarter, fairer, and more resilient for all stakeholders. The [AI for agriculture – Shaping standards for smart food systems](#) workshop provided an overview of how AI will change the digital agriculture landscape. AI encompasses a range of digital tools and platforms designed to be accessible, interoperable, and secure, fostering connectivity, communication, and innovation.

All Presentations are available [here](#).

17.1 Dialogue on digital agriculture roadmap

The dialogue discussed priorities for a collaborative Digital Agriculture Roadmap that harnesses AI and emerging technologies to build resilient, inclusive, and sustainable food systems could. Discussions explored key enablers such as data governance, interoperability, and standards, and partnerships between governments, private sector, and local communities. They also considered scalable AI based solutions to enhance productivity, reduce food waste, and empower smallholder farmers, with the aim of ensuring that digital transformation in agriculture is both equitable and future-ready. This workshop was a follow up to [FAO's First Global Dialogue on AI in Agriculture](#) held from 28 to 30 April 2025.

The development of a "Digital Agriculture & AI Innovation Roadmap" led by FAO focuses on accelerating the adoption of digital agriculture and AI to build efficient, inclusive, resilient, and sustainable agri-food systems. It aims to address global challenges such as climate shocks, geopolitical instability, resource depletion, and the demands of a fast-growing global population. Targeting governments, NGOs, academia, private firms, and investors, especially in low- and middle-income countries, the roadmap highlights advancements in AI applications, implementation challenges, and strategies for diverse agricultural settings.

Many companies in the agriculture sector still store their data in silos instead of in an integrated platform accessible to different areas of their business. For example, climate-related data can be held by a department without immediate access to data on issues such as yields and financial outcomes. Data analytics and modelling requires large amounts of data. Companies that help create and engage in data-sharing ecosystems can accumulate more data faster, benefiting the whole value chain. Integrating technologies such as remote sensing (satellites and drones), water quality sensors, and crop monitors with AI and advanced data analytics enables actionable insights for farmers. Strong data governance, meanwhile, is key to maximizing data quality and security.

An example of what is possible when data is unified would be a digital farming platform aimed at giving farmers access to a wide range of data to support regenerative efforts. The platform can obtain data from satellites, sensors, farm equipment, and other sources, and make it available via AI-driven visualization and analysis. Farmers would get insights like which biodiversity actions are best suited for their farm.

As different players invest in data collection, sharing, and governance, their efforts will benefit their partners, creating a symbiotic relationship that increases the resilience of the industry as a whole.

17.2 Demo: Project Resilience

The Irrigation Strategy part of Project Resilience from Cognizant Labs leverages the AquaCrop crop growth simulator to optimize irrigation strategies. Users can select parameters such as region and crop type, for example "Champion, Nebraska" with "Maize," to customize strategies. The demo enables adjustments of maximum irrigation (0.00–763.26 mm) and mulch levels (0.00–100%), visualizing precipitation, irrigation, yield, and mulch trends from May to October through an interactive graph.

An example scenario highlighted candidate ID "93_67," featuring a total irrigation of 407.00 mm, 4% mulch, and a final yield of 14.44 tonnes. The demo is accessible [here](#).

17.3 Farming the future – Laying the groundwork for scalable AI

Real-world AI applications from precision farming and climate forecasting to digital farmer platforms presented in this workshop, while emphasizing the importance of standards, global collaboration, and capacity-building, are shown below:

- i. IFAD highlighted AI's potential in rural development, driving climate adaptation, sustainability, and improved livelihoods. It showcased scalable pathways for food systems transformation through data interoperability, AI-driven advisory services, and climate intelligence. The "Digital Green Farmer Chat" pilot programme exemplifies inclusivity, with a majority of its users being practicing farmers and notable female participation. This programme addresses diverse farmer queries – ranging from crop management and pest control to market pricing – providing actionable and timely support. Real-world applications include climate information systems in Tanzania, credit evaluation in Mexico, and salinity alerts in Vietnam, alongside the inclusive Digital Green Farmer Chat pilot. Internally, IFAD leverages platforms like Omnidata, Blossom AI, and Harvest AI to streamline operations and enhance project design for greater impact.
- ii. Fraunhofer Heinrich Hertz Institute highlighted the transformation of agriculture through AI and IoT, moving from traditional methods to advanced digital systems. The presentation emphasized agriculture's complexity, with millions of decisions made each season, and identified key use cases and standardization gaps. Major initiatives include a digital avatar architecture to unify fragmented systems, the [NaLamKI Project](#) for an open interoperable ecosystem, and the [ACRAT project](#) promoting climate-resilient farming in Telangana, India. GeoAI and AgriFoodLoop showcase innovative applications in soil conservation and circular agriculture, with plans to scale in Argentina. The presentation also underscored the outputs of [ITU/FAO Focus Group on "Artificial Intelligence \(AI\) and Internet of Things \(IoT\) for Digital Agriculture" \(FG-AI4A\)](#) and highlighted the need for continued work beyond the focus group linking projects, data, and stakeholders to drive digital food systems forward.
- iii. Germany's Federal Ministry of Agriculture, Food and Regional Identity detailed its efforts to advance digitalization and interoperability in agriculture. The focus is on enhancing sustainability and efficiency through innovative technologies, AI applications, and knowledge transfer, supported by initiatives like the "Digital Trial Fields." While 5G infrastructure is expanding, challenges remain in achieving seamless data exchange between devices, prompting over €15 million in funding for projects addressing open-source software and coordinated machine management.
- iv. Abris Inc. shared insights on "Letzfarm: AI-Driven Smart Farming Education for Smallholder Farmers" showcasing how AI can empower smallholder farmers to achieve

sustainable agriculture, higher yields, and greater resilience against climate challenges, particularly in vulnerable small island economies. The mobile-first Letzfarm platform delivers personalized recommendations by integrating weather patterns, soil conditions, pest management, post-harvest storage, and real-time market prices, resulting in measurable impacts such as a 15% yield increase and 20% reduction in crop losses for 2,000 farmers in Trinidad and Tobago. LetzFarm is expanding through global data protocols, regulatory frameworks, localized training for Small Island Developing States (SIDS), and educational partnerships with universities for AI-based agricultural training and research.

- v. WFP highlighted how AI can reshape food systems by shifting from reactive crisis response to proactive risk management, enabling rural communities and national planners to anticipate shocks like droughts, floods, and market volatility. WFP partnerships with Google, Microsoft, IBM, and the Gates Foundation have advanced food insecurity forecasting, crop yield prediction, micronutrient analysis, and geographic targeting. A key focus is on strengthening data channels, real-time analytics, and scalable AI models, with successful pilots in Yemen, Nigeria, Cameroon, and Ethiopia. While AI is powerful, its impact depends on complementing human expertise and being built on trust, transparency, and ethics to enable faster, more targeted interventions for food security.
- vi. The German Agency for International Cooperation (GIZ) and the Gates Foundation focused on the challenges faced by smallholder farmers, including limited access to personalized digital advisory services, and the lack of solutions designed for low-literacy, low-digital skill groups. They outlined the opportunity to leverage recent advancements in AI and digital technologies to deliver personalized and dynamic information to farmers. A use case was presented where a farmer named Rose uses a feature phone and Interactive Voice Response to receive farming advice in her local dialect. The project timeline includes key milestones from 2023 to 2025, such as the call for proposals, Minimum Viable Product development, and end-user research and evaluation. Lessons learned include the need for multi-disciplinary partnerships and further investments in foundational technologies to sustain agricultural advisory systems.

17.4 Key outcomes

Key outcomes and findings can be summarized as follows:

1. There is a need for transparent AI standards in digital food, with ITU positioned to collaborate with FAO and other UN partners on global frameworks for data governance and AI governance toolkits.
2. Gaps were identified with respect to open data protocols, and the development of scalable digital advisory platforms, particularly for smallholder farmers with low literacy and limited digital skills.

3. The Digital Agriculture & AI Innovation Roadmap and NaLamKI project highlight opportunities for ITU to support the standardization of AI interfaces, data-sharing protocols, and validation of AI-powered digital food systems.
4. Multi-disciplinary partnerships and inclusive solutions will be key for standards to evolve alongside emerging technologies for resilient and future-ready food systems.
5. During the AI for Good Global Summit 2025, the [Global Initiative on AI for Food Systems](#) was launched. It is led by ITU, FAO, WFP, and IFAD. It aims to drive shared digital infrastructure, pilot projects, and standards to empower governments and innovators to deliver real-world impact for resilient food systems. The outputs of this initiative will feed into key ITU standardization workstreams to support adoption at scale.

18 Women leaders in AI and standards

This event was jointly organized by the [Network of Women in ITU-T](#) (NoW in ITU-T) and [Women in AI](#). Women leaders and AI experts explored the transformative impact of increasing women's participation and leadership in AI and standards development, from driving innovation and boosting revenue to enabling inclusive growth in the technology sector. Through compelling success stories and evidence-based insights, the session presented recommendations and a call to action for companies and governments to advance women's leadership in AI and standards.



Figure 50: Doreen Bogdan-Martin, Secretary General, ITU delivering her opening remarks at the event.

Key points made by the various speakers are summarized below:

- The key findings of the ITU-T survey on "Barriers to women's participation in standards" and the recommended actions for the ITU secretariat and ITU membership to increase engagement, were shared.
- Action to promote gender-responsive policy is important and the AI for Good Global Summit could serve as an advocacy platform.
- Gender diversity in AI standardization work is a business imperative.
- A call was made for more initiatives to create pipelines encouraging more girls to enter STEM fields.
- Female participation in the ITU AI/ML Challenges remains low but is gradually increasing.
- Youth visibility on panels can be a strong inclusion tool and speakers encouraged co-designing AI –governance frameworks and standards with young people.



Figure 51: Left to right: Paola Cecchi Dimeglio, Chair, Executive Leadership Research Initiative for Women and Minority (ELRIWMA), Professor, Harvard University; H.E. Ms. Neema Lugangira, Women Political Leaders' Secretary General, Member of Parliament, Tanzania (United Republic of); Alessandra Sala, Global President of Women in AI; Chair of AI and Multimedia Standards Collaboration, Sr. Director of Artificial Intelligence and Data Science, Shutterstock; Salma Abbasi, Founder, Chairperson and CEO, eWorldwide Group; Kathleen A. Kramer, President & CEO, IEEE; Celina Lee, Co-Founder & CEO, Zindi Africa; Roser Almenar, ITU Secretary-General Youth Advisory Board member, PhD Candidate, AI & Space Law at the University of Valencia

Recognizing that more women in AI and standardization will drive economic value, speakers called for practical STEM pipelines, youth codesign, strong networks, and targeted actions at the institutional level, alongside increased funding to support measurable participation outcomes.

The session concluded with a Call to Action emphasizing the importance of networks to advance women's careers in AI and standardization.

Call to Action:

“We call on governments and companies to act with intention and foresight: to dismantle systemic barriers, invest meaningfully in the advancement of women, and embed inclusive leadership in shaping the future of AI and international standards. The future is being written now, let's write it together.”

Part 3: Future AI standards for frontier technologies

19 AI and virtual worlds: Building the cities and governments of tomorrow

AI and virtual worlds are redefining governance, infrastructure, and public services. The summit's workshop on the topic welcomed representatives of governments, cities, and international organizations to align national and local strategies with global priorities outlined by the Pact for the Future and its Global Digital Compact. The message was clear: AI-powered virtual worlds must serve the public good and should be anchored in shared values and supported by international standards. The event also highlighted the growing global backing for the [Global Initiative on Virtual Worlds and AI-Discovering the Citiverse](#). Launched by ITU, the UN International Computing Centre (UNICC), and Digital Dubai, the initiative is now supported by over 60 partners and serving as a springboard for collaborative policymaking, innovation exchange, and standardization efforts aligned with the Pact for the Future and its Global Digital Compact. In addition, the event invited all stakeholders to the [3rd UN Virtual Worlds Day](#), set for 11–12 May 2026, to celebrate impact, share results, and align global efforts.

19.1 Virtual worlds, metaverse and citiverse – need for standards

The [ITU Focus Group on Metaverse](#) laid the groundwork for new ITU standards for virtual worlds. The "citiverse" concept, a key area of study for the focus group, leverages a combination of immersive, intelligent, and interconnected digital technologies to support cities in adopting a human-centred approach and advancing sustainable development. It is designed specifically for urban contexts, enabling the creation of interconnected digital twins of cities, where inhabitants, businesses, and governments can collaborate to address complex urban challenges and build more sustainable, inclusive, and resilient communities. The citiverse could offer new administrative, economic, social, cultural, and policy-making capabilities, providing virtual goods, services, and participatory tools to city and community actors – such as citizens (represented as digital avatars), local authorities, service providers, and civil society organizations.

Such AI-powered virtual worlds must be built on a foundation of trust, inclusion, and interoperability. [ITU-T Study Group 20 \(IoT, digital twins and smart cities\)](#) serves as ITU's lead expert group for the development of international standards for the citiverse.

The [Global Initiative on Virtual Worlds and AI – Discovering the Citiverse](#) is a global multistakeholder platform to foster open, interoperable, and human-centric virtual worlds through three strategic pillars: Guidance, implementation, and global governance.

Key Pillars of Action

The Initiative advances its mission through three strategic pillars, each supported by dedicated tracks that tackle the most pressing opportunities and challenges in AI-powered virtual worlds. These pillars represent a comprehensive and inclusive approach—from developing global guidelines to driving real-world implementation in cities around the world.



Figure 52: Global Initiative on Virtual Worlds and AI -Key Pillars of Action

Box 4: Why international standards and guidelines matter for the civerse and AI-powered virtual worlds

International standards and guidelines can support trusted civerse and virtual worlds. They enable interoperability between platforms and systems, helping cities and communities to connect seamlessly across verticals. They promote accessibility by design for all inhabitants – including persons with disabilities and those in underserved areas – to benefit from AI-powered virtual services.

In addition, international standards provide a common framework to align digital infrastructure, urban innovation, and public services with global best practices. This not only helps to avoid fragmentation but also accelerates deployment, reduces costs, and supports sustainable development at scale.

[Study Group 20](#) is developing several international standards to guide civerse deployments, including:

- [Y.civerse-reqts – Requirements of civerse platform for smart sustainable cities and communities](#)

Global Initiative on Virtual Worlds and AI - Discovering the Civerse Civerse Use Case Taxonomy Overview Use Case Identification Track itu.int/metaverse/virtual-worlds/ Civerse Initiative	Complementing this work, the Global Initiative on Virtual Worlds and AI is advancing guidelines for the civerse. One of its key deliverables is: • Civerse Use Case Taxonomy Overview
---	--

The Global Initiative aims to shape a future where AI-powered virtual worlds are inclusive, trusted, and interoperable. By connecting people, cities, and technologies, it empowers meaningful progress through AI-powered virtual worlds. Urban innovation is key in shaping

new digital public infrastructure. Tools like the [Citiverse Use Case Taxonomy](#), developed by the Global Initiative, offer practical guidance for aligning AI-powered virtual world deployments with urban priorities with the aim of ensuring inclusive, trusted, and interoperable digital ecosystems.

AI-powered virtual worlds, including the metaverse and the citiverse, require common standards to be interoperable, secure, inclusive, and trustworthy. As AI increasingly drives how people, cities, and communities interact in digital environments, international standards can help enable seamless connectivity between platforms, safeguard users, and support human-centred, sustainable development in both physical and digital spaces.

19.2 Policy frameworks for virtual worlds

Robust policy frameworks can help ensure that virtual worlds – often powered by AI and integrating digital twins, IoT, and immersive technologies – are safe, inclusive, ethical, and aligned with UN Sustainable Development Goals and the Global Digital Compact.

Appropriate policy frameworks can help governments and communities harness virtual worlds responsibly while protecting users and fostering innovation. Virtual worlds have evolved far beyond entertainment and can become powerful drivers of cultural and economic value on a global scale. For example, the Dominican Republic has launched the "Gobverse," an immersive government platform that uses virtual worlds and digital twins to modernize public services and leapfrog traditional infrastructure challenges. Similarly, the AI Hub for Sustainable Development, launched under the Italian G7 Presidency, represents a bold step toward bridging compute gaps with and for Africa.

To unlock real-world impact, policy priorities could include:

- a) Fostering collaboration among governments, technology providers, academia, and civil society to shape inclusive and future-ready cities and communities.
- b) Driving economic growth through AI-enabled innovation ecosystems that create new jobs, smarter products, and sustainable business models.
- c) Scaling global AI governance through adaptable, locally grounded frameworks that empower the Global South to harness these technologies to address pressing development challenges.

20 Brain-computer interfaces

As a frontier intersection of neuroscience, AI, and engineering, Brain-Computer Interface (BCI) technology is transforming human-machine interaction by enabling direct communication between the brain and external devices, bypassing peripheral nerves and muscles. This breakthrough unlocks new possibilities across healthcare (e.g. restoring motor functions for paralyzed patients via neural signal decoding), industrial control (e.g. brain-controlled drones navigating complex terrains for emergency rescue), and daily life applications, while addressing longstanding accessibility barriers for persons with disabilities.

The workshop gathered technology experts, policymakers, and business leaders to explore BCI's latest advancements, including notable progress in both invasive technologies (such as long-term stable neural implants) and non-invasive solutions (like portable, user-friendly electroencephalogram (EEG) headsets).

Discussions focused on key challenges that hinder widespread adoption, including:

- Technical bottlenecks in signal accuracy and stability (especially for non-invasive systems in complex tasks).
- Ethical risks around neural data privacy and potential misuse.
- Fragmented technical standards that limit cross-border collaboration and scalability.

The goal was to identify actionable pathways for standardization, ethical governance frameworks, and global cooperation, with the aim of ensuring that BCI technology develops in a way that is safe, inclusive, and aligned with human-centric values to maximize its societal impact.

Presentations could be found at the [BCI Workshop](#) website.

20.1 Keynote address

ITU showcased BCI's life-changing impact through case studies including a Portuguese locked-in syndrome patient communicating via BCI and 5G, and a Brazilian patient controlling a car with thought. Some of the key challenges highlighted were upgrading signal acquisition safety, addressing EEG data scarcity, and strengthening ethical governance. ITU's collaboration with WHO on medical BCI guidelines was highlighted as a step towards global norms.

University of Bath and Fudan University provided an overview of BCI's industrialization barriers, noting the lack of a mature "basic research-clinical development-approval" pipeline compared to pharmaceuticals. There are some success stories however, such as FDA/NMPA-approved stroke rehabilitation devices covered by medical insurance, stressing that standardization (e.g., EEG recording guidelines) is critical for scaling.

China's Northwestern Polytechnical University presented breakthroughs in brain-controlled systems, including an 86.5% success rate for drone swarms navigating obstacles via motor imagery signals. Challenges outlined included non-invasive signal noise (accuracy <90% for complex tasks), 300ms response delays, and individual EEG variability. Solutions included multimodal data fusion and adaptive algorithms, with international partnerships driving open-source platforms.

20.2 Opportunities and challenges of BCI

The following opportunities and challenges of BCI were identified in this session based on the presentations made:

- a) University of Rome showcased BCI's application in monitoring high-risk professionals like pilots and surgeons. By analyzing EEG, galvanic skin response, and heart rate via machine learning, the system real-time quantified cognitive load, attention, and stress to trigger safety alerts. Scalable lightweight devices are required—using few electrodes while retaining 80-90% core data—for easy real-world deployment.
- b) Neuroscience-Paris Seine-IBPS Laboratory ICNRS INSERM at Sorbonne University warned of ethical risks in non-medical BCI expansion, particularly in employee monitoring scenarios that could infringe on cognitive freedom. It was highlighted that

there is a need for robust governance frameworks, citing the OECD's 2019 recommendations (classifying neural data as uniquely sensitive) and France's 2022 Responsible Development of Neurotechnology Charter – a multistakeholder agreement banning non-voluntary neural data use and protecting mental privacy, now integrated into the European Brain Initiative.

- c) A BCI implant recipient shared firsthand experience of controlling prosthetics with four implanted electrodes, highlighting needs for longer device lifespan (current maximum is 10 years) and easier maintenance.
- d) Institute of Information and Communication Technologies (IICT) at the Bulgarian Academy of Sciences focused on technical challenges of invasive BCI, such as signal instability caused by brain deformation. The need for standardized failure analysis frameworks was emphasized and the presentation also highlighted IEEE's efforts to unify terminology, data formats, and metadata compatibility to break resource-sharing barriers.
- e) China's Beijing Institute of Technology analysed BCI's market growth and risks, including surgical infections from invasive devices which have a 5-10 year implant lifespan, and ethical issues like neural data privacy leaks. The presentation highlighted global governance progress, such as China's BCI Research Ethics Guidelines and UNESCO's 2025 draft Recommendation on the Ethics of Neurotechnology.

20.3 Industrialization and commercialization barriers of BCI

The following industrialization and commercialization barriers of BCI were identified in this session based on the presentations made:

- a) Araya showcased 79.8% accuracy in decoding 512 common phrases from EEG data, enabled by a "health data pre-training + patient fine-tuning" framework that slashed clinical adaptation time – boosting word recognition accuracy from 13.2% (using only patient data) to 54.5% with minimal patient-specific data. It was highlighted that a rental-based business model (equipment leasing + pay-per-service) to lower initial hospital investment barriers, is now piloted in three Japanese rehabilitation facilities to accumulate real-world data.
- b) SceneRay presented dual-target deep brain stimulation (Combo-Stim DBS) therapy for opioid addiction, which precisely targets the nucleus accumbens and anterior cingulate bundle to repair reward circuit deficits. Clinical trials (2018–2024) and 10-year follow-ups showed an 80% 6-month abstinence rate and 69% long-term success, with no adverse impacts on cognition or daily function. Adhering to ISO 13485 (quality management) and ISO 14971 (risk management) standards, the technology supports safe device removal after two years, paving the way for scalable medical applications in addiction treatment.
- c) iFLYTEK European Division shared their experience in scaling neurotechnology applications through "R&D-scenario validation-large-scale promotion" models. For example, optimizing algorithms for specific scenarios and reducing deployment costs, which offers insights for BCI's "technology inclusivity" challenges.

20.4 Standardization and global cooperation

The following areas where standards are needed and where global cooperation can help were highlighted:

- a) Northwestern University stressed AI's role in BCI standardization, focusing on objective efficacy metrics derived from biomarkers. Through research by the post-traumatic stress disorder (PTSD) team, they identified α -wave rebounds (a measurable neural signature post-treatment) as a standardized indicator, correlating with 36% symptom reduction. These metrics support cross-border clinical collaboration by enabling consistent evaluation, while integrating cultural adaptability in research designs to enhance global standard relevance.
- b) CAICT, focused on BCI standardization needs and global collaboration pathways. The BCI industry faces core challenges such as poor device interoperability, unclear safety benchmarks, and fragmented neural data governance. Standards are needed to address these challenges. The presentation highlighted progress in China, such as the group standards for EEG-based attention monitoring systems. Emphasizing the role of international platforms like ITU, the presentation called for global synergy to align technical concepts and testing methods for standards to support safe and inclusive BCI development across regions.
- c) University of Bath & Fudan University emphasized that standardization should act as a "universal interface" rather than restricting innovation. Standards should clarify core performance indicators (e.g. command recognition delay and long-term stability) while leaving room for diverse technical routes (e.g. invasive vs. non-invasive devices). For instance, unified EEG data interaction protocols could enable cross-institutional collaboration in epilepsy monitoring research without limiting hardware innovation.

20.5 Key outcomes

The key outcomes for the workshop can be summarized below:

1. **Technical progress:** BCI has advanced in both invasive (e.g. long-term implant stability) and non-invasive (e.g. lightweight EEG headsets) technologies, with applications spanning healthcare, industry, and safety.
2. **Industrial barriers:** Scaling is hindered by signal accuracy issues, high costs, fragmented standards, and inadequate clinical pipelines. Cross-sector collaboration for interoperability and standards (hardware + algorithms + clinics) are critical.
3. **Ethics and privacy:** Neural data sensitivity demands global frameworks to prevent misuse, with emphasis on consent and transparency.
4. **Standardization and global collaboration:** Areas of focus for ITU standardization could include unified standards for data formats, performance metrics, accessibility, network infrastructure, and safety specifications to enable interoperability and trust.

In conclusion, BCI technology is at a pivotal stage, with the potential to revolutionize human-computer interaction and improve lives globally. Advancing this technology requires

technological innovation, ethical safeguards, scenario-driven standards, and global collaboration to ensure inclusive and sustainable development.

21 Building the technical foundations for embodied intelligence in connected ICT environments

This workshop addressed embodied intelligence and its dependence on advanced ICT infrastructures, with the goal of identifying standardization priorities. Panelists included representatives from academia, industry, and standards bodies. The discussion followed four themes:

- i. Definitions
- ii. Technical challenges and use cases
- iii. Standardization gaps
- iv. Pathways for collaboration

21.1 Defining embodied intelligence

The session opened by clarifying the concept of embodied intelligence. A definition from the Chinese robotics industry community describes it as intelligent systems interacting with the environment through physical entities such as robots, capable of environmental perception, cognition, autonomous decision-making, and action execution.

Three main distinctions from traditional AI were outlined:

- Physical dependence: Embodied AI operates through tangible hardware with physical constraints (e.g. robotic arms, sensors, and motors)
- Perception-action loop: Continuous real-time decision-making under physical laws, requiring feedback and adaptive strategy.
- Irreversible consequences: Actions have direct physical effects, unlike reversible digital outputs.

Additional interventions emphasized that embodied systems are inherently social when operating in human environments. The form, movement, and appearance of robots influence trust, usability, and human expectations. It was noted that robots, even when not designed for social interaction, often elicit social responses due to their presence in human spaces. Attention was called to the importance of sensitivity to context and awareness of the risk of embedding stereotypes into robotic embodiments.

21.2 Technical and infrastructural challenges

Industry participants highlighted latency, connectivity, and bandwidth as critical barriers. Teleoperation across continents was reported to create delays of up to 2-3 seconds, undermining usability. Stable VR-based teleoperation was said to require bandwidths of around 100 Mbps.

Further challenges discussed included:

- Accessibility: Dependence on cloud infrastructure limits use in rural and underserved regions, often where robots are most needed.

- Privacy and control: Users risk losing functionality when services close (e.g. the "Boxy" children's robot that stopped working when its company shut down). Concerns were raised about where data goes, who has access, and how secure it is.
- Security: Robots and their telecom interfaces can be "data leaky" and hackable, exposing biometric data such as eye tracking, facial expressions, and haptic signals. Risks are significant in sensitive contexts like healthcare or elderly care.

21.3 Standardization gaps and needs

Speakers identified priority areas where standards are needed:

- Definitions and taxonomy
- Component-level clarity on building blocks and interoperability
- Simulation-to-reality gap: Bridging synthetic training data and real-world environments.
- Benchmarks: Application-oriented benchmarks are needed to reflect real-world requirements.
- Socio-technical standards: In addition to technical specifications, standards development should consider socio-technical dimensions such as human rights,

Standardization activities of ITU-T on embodied AI is shown in Box 5 below, including standards under development and published.

21.4 Collaboration across standards bodies

Participants agreed that international collaboration is essential and made for:

- Multistakeholder inclusion, with startups, emerging companies, academia, and civil society engaging alongside large institutions.
- Parallel pre-standardization discussions supporting formal standards-development processes.
- The consideration of human rights, privacy, dignity, children's rights, and accessibility alongside technical discussions.
- accessibility, and environmental impact.

Box 5: ITU standardization activities on embodied AI

ITU-T Study Groups are actively engaged in standardization work on and related to embodied AI. Information about related published standards and ongoing work items are provided below.

ITU-T Study Group 2

ITU-T Study Group 2 has published two ITU-T Recommendations related to embodied AI.

- **Recommendation [ITU-T M.3167.1 \(03.2025\)](#)**
With the continuous development of Internet of things (IoT) technology, the application of intelligent maintenance robots (IMRs) in the field of telecommunication smart maintenance (TSM) is increasing. Recommendation ITU-T M.3167.1 provides the requirements for the interface between the IMR-based smart patrol system (IbSPS) and the telecommunication smart maintenance system (TSMS) at a protocol-neutral level. It describes the position of the relevant interface and specifies the high-level requirements for interface interaction, as well as specification level use cases for each requirement.
- **Recommendation [ITU-T M.3367 \(04/2023\)](#)**
Recommendation ITU-T M.3367 introduces requirements for intelligent maintenance robots (IMRs) based on-site smart patrol of telecommunication networks, includes the network elements to be patrolled, requirements for management function of IMR-based patrol and related management interfaces.

ITU-T Study Group 11

ITU-T Study Group 11 is actively working on intelligent edge computing and its related use in embodied AI systems. Within this context, SG11 developed two ITU-T Recommendations

- **Recommendation [ITU-T Q.5029](#)**: "Data management interfaces in digital twin smart aquaculture system with intelligent edge computing".
- **Recommendation [ITU-T Q.5030](#)**: "Data management interfaces for intelligent edge computing-based flowing-water smart aquaculture system."

ITU-T Study Group 13

ITU-T Study Group 13 (SG13) has addressed "Robotics as a Service" (RaaS) in its work, specifically through **Recommendation [ITU-T Y.3533](#)** which outlines the functional requirements for developing, operating, and managing robots within a cloud computing environment.

ITU-T Study Group 20

Several ITU-T Study Group 20 Recommendations (such as ITU-T Y.4106, Y.4607, and Y.4506) directly address embodied AI systems, defining requirements and architectures for robots operating in real environments.

Recommendations [ITU-T Y.4471](#) and [ITU-T Y.4487](#) provide supporting network and sensor infrastructure that enables embodied AI systems, such as autonomous vehicles, to function safely and effectively.

ITU-T Recommendations [Y.4607](#) and [Y.4506](#) provide the requirements and reference architecture for autonomous delivery robots — a concrete application of embodied AI in urban environments.

ITU-T Study Group 21

[ITU-T SG21](#) is actively developing standards that shape the foundations and evaluation of **Embodied Artificial Intelligence (EAI)**.

Two following core work items directly target EAI by defining system capabilities (perception, decision-making, action, and interaction) and establishing standardized benchmarking frameworks:

- **[E.RF-EAI](#)** (*Requirements and framework for embodied artificial intelligence systems*);
- **[E.EAI-AC-BK](#)** (*Assessment criteria for Embodied Artificial Intelligence system: Benchmark*).

Box 6: ITU standardization activities on embodied AI (continued)

ITU-T Study Group 21 (continued)

The following approved ITU-T Recommendation contributes to embodied AI interaction scenarios:

- [F.746.9](#) (*Requirements and architecture for indoor conversational robot system*).

The following complementary efforts address enabling technologies. These introduce evaluation methods for foundation models, including aspects of visual embodied intelligence and multimodal reasoning critical to Embodied AI:

- [H.FDM-AC-Gen](#) (*Assessment criteria for foundation models: General*);
- [H.FDM-AC-V](#) (*Assessment criteria for foundation models: Vision*).

Domain-specific applications are being studied in F.MTTIR, which operationalizes embodied perception-action loops in inspection contexts for safety and efficiency:

- [F.MTTIR](#) (*Requirements and framework of computer vision and audition-based transportation tunnel inspection robotic systems*).

Additionally, SG21 is advancing ontology-based standards to enable context modelling and interaction in intelligent spaces, supporting adaptive and interactive EAI deployments:

- [F.DCOR-reqs](#): (*Requirements for construction of ontologies knowledge base for household service robots in intelligent spaces*)
- [F.DCHE-Req](#)s: (*Requirements for construction of ontology knowledge base of home environment for robot interaction in intelligent space*)
- [F.DCHU](#): (*Requirements for construction of home user profile ontology knowledge base for robot interaction in intelligent space*)

Collectively, these initiatives position SG21 as a key driver in defining **requirements, benchmarks, and enabling ecosystems** for embodied intelligence.

21.5 Looking forward

Participants concluded that embodied AI is still at a very early stage of deployment. Priorities identified for the coming year include:

- Continued studies on embodied in AI in ITU-T Study Group 21 embodied AI.
- Engaging broader communities of potential users, including youth.
- Establishing benchmarks and performance metrics that reflect real-world applications.
- Developing interoperability frameworks for robotics in urban environments and across connected devices.
- Exploring mechanisms for trust and reliability in anticipation of the emergence of billions of future AI agents.

The session closed with calls for collaboration, inclusivity, and coordination across technical and socio-technical domains.

Part 4: Quantum for Good

22 Quantum for Good Track

Quantum technologies are poised to revolutionize many industries, solving complex problems that until now seemed unfathomable. Imagine the possibilities for drug discovery, disease risk prediction, and biological research in high-impact areas like protein folding, as well as in weather forecasting and climate modelling.

Quantum technologies are advancing rapidly, beyond lab experiments toward marketplace applications, promising to transform sectors from healthcare and climate science to cybersecurity and communications. Realizing this potential, however, their impact must be inclusive, ethical, and sustainable. To achieve this, the Quantum for Good track of the AI for Good programme was launched at the 2025 summit with the aim of driving responsible innovation, strengthening global collaboration, and supporting the development of inclusive standards to ensure quantum technology delivers real-world benefits for all.

Framed within the UN's International Year of Quantum (IYQ 2025), Quantum for Good discussions at the 2025 summit examined how innovation in the field can serve the public good. The discussions focused on governance, equitable access, and workforce readiness, while highlighting urgent needs for security, standardization, and ethical guidance.



Figure 53: Left: Cierra Choucair, CEO, Universum Labs, Strategic Content Director, The Quantum Insider. Right: Qiang Zhang, Executive director, Jinan Institute of Quantum Technology, Professor, University of Science and Technology of China

Relevant sessions at the summit showcased how quantum science and technology complements AI in tackling global challenges and provided insights on the following aspects:

- a) The maturity of quantum technologies on a 20-year timeline, from near-term sensing devices to long-term visions of fault-tolerant quantum computing.
- b) How quantum is transforming sectors
- c) Trust and post-quantum security

Beyond dialogue, the track also featured an artistic showcase from Mekena McGrew, Quantum musician and Wiktor Mazin, Quantum fractal artist, who premiered their collaborative work "Echoes from the Quantum", a performance which blended synthesized quantum music, generative quantum fractal art, and patterns inspired by quantum algorithms.

Drawing on the metaphor of quantum entanglement, the performance highlighted how global scientific progress depends on shared understanding and cooperation across borders. By transforming complex quantum phenomena into sound and visuals, Echoes from the Quantum offered the audience not only a glimpse of the beauty within quantum systems, but also an invitation to experience the spirit of collaboration that drives the field forward.

Full session descriptions are available through the links below:

- a) [Quantum for Good: Industry leadership, innovation and real-world impact](#)
- b) [Quantum technology and diplomacy: Why should we care?](#)
- c) [Quantum for all: Innovation, access & impact](#)
- d) [Building tomorrow's quantum workforce: Voices of future leaders in quantum](#)
- e) [Quantum for Good: Shaping the future of quantum – What happens next?](#)

22.1 From promise to progress: Focusing on real use-cases

Quantum technologies are still at a nascent stage, calling for collaborative, inclusive, and purpose-driven action that links research to real-world impact and global development goals.

A key insight was distinguishing between what is “*technically ready*” from what is truly relevant. In a rapid-fire kick-off panel, experts from quantum computing, communication, and sensing fields placed themselves on a 20-year timeline (from near-term to long-term) to indicate how close their fields are to maturity. The exercise yielded a visual consensus: quantum sensing applications appear closest to real-world impact, quantum communication and networking are mid-term, and fault-tolerant quantum computing remains on the far horizon. The point was not to pinpoint exact dates, but to recognize that each area faces different hurdles and timelines. The session underscored that there is no single quantum timeline; instead, multiple intersecting timelines are unfolding, each shaped by technical challenges and practical use cases.



Figure 54:Seizo Onoe, Director of the Telecommunication Standardization Bureau (TSB), International Telecommunication Union (ITU)

“We should take every opportunity to share information and help position everyone for success. This new Quantum for Good track will help us do exactly that.”

Seizo Onoe

Panelists agreed on the need for quantum research and development to remain problem-driven, collaborative, and guided by real-world impact. Their perspectives also reinforced that, across all areas of quantum technology, use cases need to be prioritized based on a clear assessment of their real-world impact. This approach helps ensure that progress is collaborative, problem-driven, and focused on meaningful outcomes rather than hype.

Momentum is shifting from experimentation toward actual deployment as quantum technology solutions are beginning to tackle real-world problems. One session highlighted examples of use cases being deployed on the ground, focusing on practical applications and the underlying principles enabling them. Panelists presented how quantum sensing is being piloted in water resource management to optimize usage, how quantum algorithms aid climate adaptation strategies, and how quantum cryptography is strengthening secure communications infrastructure.

The first real-world applications of quantum technologies for social good are beginning to emerge and the community is learning to prioritize impact and concentrating on "early win" applications where quantum technologies add value. However, to scale, the field requires equitable access, cross-domain collaboration, and clear application development guidelines.

22.2 Quantum dilemma and risks

The advent of quantum computers brings not only hope, but also risks; especially to the security of our digital systems. One urgent topic discussed was the looming threat that quantum computing poses to current cryptography. Today's widely used public-key encryption methods (like RSA and ECC) rely on mathematical problems that a sufficiently advanced quantum computer could solve exponentially faster, potentially decrypting sensitive data that is now considered secure.

Experts warned that this "Q-day" (when quantum code-breaking becomes feasible) appears to be drawing closer, perhaps sooner than previously assumed. In fact, one cybersecurity panelist cautioned, "We are already too late. The quantum threat timeline has accelerated from a decade away to much sooner." This stark warning underlined the need for immediate action to protect information in the quantum era as malicious actors may already be harvesting encrypted data today, hoping to decrypt it later, so the clock is ticking for organizations to migrate to quantum-safe solutions.

The experts urged a multi-pronged response:

1. **Map critical assets:** Organizations should first inventory their systems and data to understand what sensitive information might be vulnerable to future quantum decryption.
2. **Educate stakeholders:** From C-suite executives to IT personnel, all stakeholders need to understand the risk and the urgency of transitioning to new cryptographic standards
3. **Act early:** Organizations should begin integrating quantum-resistant encryption algorithms and protocols into their infrastructure well before large-scale quantum computers become reality

International standards bodies and national institutes have already been vetting candidate algorithms for post-quantum cryptography; the challenge now is deploying them in time. The clear consensus was that the world should not wait for a definitive "quantum computer moment" to act. Every additional year of delay in rolling out quantum-safe encryption increases the chance that secure communications, or stored data could eventually be compromised. "Quantum-proofing" cybersecurity is an urgent priority that requires forward-thinking efforts today, not tomorrow.

22.3 Inclusion and workforce: Building capacity for all

Ensuring that the quantum revolution benefits everyone was another key focus of the discussions. Participants stressed that realizing "quantum for good" requires constantly asking "for who's good?" and taking deliberate steps to make quantum technology inclusive. This starts with broadening access to education and resources in this highly specialized field. As quantum science advances, there is a risk that only some institutions and well-funded companies in a few countries will cultivate the needed expertise, leaving others behind. To counter this, the community is prioritizing efforts to democratize quantum knowledge and nurture talent globally.

Roughly one-third of the world's population (around 2.6 billion people) still lack Internet access, which is a stark reminder that without basic connectivity and digital infrastructure, many developing regions would not be able to participate in cutting-edge technology fields like

quantum technology. Bridging this digital divide is a foundational step toward quantum inclusion.

Growing and diversifying the quantum workforce is equally urgent. In a forward-looking panel, young professionals and "future leaders" in quantum described how entering the field can feel like a Catch-22 where "experience is needed to get experience." Breaking this cycle will require lowering barriers to entry by creating more early-career opportunities such as internships, apprenticeships, and hands-on training programmes, enabling students and recent graduates to gain practical exposure even without advanced degrees. They also highlighted the importance of fostering a supportive community and mentorship for newcomers, rather than the field becoming a field of exclusivity.

The skillset demanded by the quantum workforce is also expanding. Beyond physicists and theorists, the industry will need a range of roles including engineers, software developers, system integrators, technicians, as well as experts in ethics, law, and communications. Cross-disciplinary skills were highlighted as vital, for example training physics students in communication and teamwork so they can collaborate effectively and explain complex ideas to non-experts. Public outreach and education are essential to gain public trust and inspire the next generation. Practical suggestions included integrating communication and ethics modules into technical curricula and engaging schools and communities to raise quantum literacy.

The overarching message was that building "quantum for all" requires intentional effort now. This means extending resources and lab access beyond established centres of excellence, supporting women and under-represented groups in pursuing the field, and sharing educational content widely. International cooperation in workforce development was called for to share best practices and even create joint programmes to train quantum engineers and scientists.

The discussions called for global resolve to help countries grow their expertise, build a quantum-ready workforce, create opportunities (including for women and girls), equip youth for further quantum development, and foster quantum research locally. Ensuring diverse participation is not only a matter of fairness but a driver of creativity and innovation. A broader pool of participants means a wider range of ideas for how quantum can be applied to address global challenges and deliver real benefits to humanity.

22.4 International collaboration and standards for quantum

The rapid development of quantum technologies is outpacing policy development, raising urgent questions around governance, security, and equitable access. The dynamic and global nature of quantum innovation highlights the value of joint research, open innovation, and international consensus-based standards for quantum technologies to deliver global benefits while guarding against related risks.

A central concern raised in the discussions was the risk of the "quantum divide." Today only a handful of countries are making major investments in quantum research and development, positioning themselves to leap ahead while many others lack even the digital infrastructure to participate. Without deliberate action, this pattern could mirror or exceed the digital divide of the Internet era. Strong diplomatic and multilateral measures, capacity-building programmes, knowledge transfer, and inclusive research partnerships are critical to help developing regions build expertise and access quantum benefits.

Leaders from QED-C, QuIC, and Q-STAR highlighted how international quantum consortia navigate collaboration in a geopolitically fragmented world. While acknowledging national security concerns and diverse organizational structures, there was broad consensus on the need for cross-border efforts and proactive consideration of both benefits and dual-use risks. Challenges include differing national priorities, rapid technological evolution, error correction and scalability issues, and balancing collaboration with export controls.

Standards are one of the most practical and agile forms of collaboration. Technical standards can help harmonize efforts across borders, create a shared vocabulary, and support the interoperability and security of emerging quantum systems. Standards also provide a framework for building trust, reducing ambiguity, and potentially laying a foundation for upcoming regulation, without stifling innovation during the nascent stages of technology development.

ITU standards are addressing network and security aspects of quantum information technologies with an initial focus on Quantum Key Distribution (QKD) for quantum-safe encryption and authentication. This work has involved over 300 experts from 180 organizations and 43 countries. [ITU-T Y.3800](#) “*Overview on networks supporting quantum key distribution*” provides the foundational framework for quantum key distribution (QKD) networks, covering architecture, terminology, and design principles. [ITU-T X.1710](#) “*Security framework for quantum key distribution networks*” builds on ITU-T Y.3800, identifying security threats, and specifying security requirements and measures for QKD networks. [ITU-T X.1811](#) focuses on 5G, assessing quantum threats to existing cryptography and providing guidance for applying quantum-safe algorithms. Together, these and other standards in the ITU-T suite give governments and industry a clear reference for building secure, interoperable quantum-safe networks. [Ongoing ITU-T work](#) also addresses quality assurance, protocols, and even the use of machine learning in QKD networks, laying the foundation for a scalable, globally trusted quantum infrastructure.

In the current phase where applications are largely pre-commercial and risks still theoretical, global standards can offer consistent technical language, testing protocols, interoperability guidelines, and benchmarks. These mechanisms can help governments, industries, and researchers align efforts, prevent fragmentation, and coordinate around shared goals.

Beyond technical aspects, international standards guide ethical practice and responsible innovation. Participants highlighted the dual-use nature of quantum technologies; capable of strengthening climate forecasting or healthcare on one hand, but also of breaking encryption or enabling new weapons on the other. Drawing lessons from AI, they urged the community to embed ethics, security, and equity considerations from the outset, rather than as an afterthought. Multi-stakeholder spaces where scientists, policymakers, and industry leaders could convene to co-design guidelines and ethics were seen as vital for addressing these dilemmas before competitive pressures harden.

From establishing international standards and security frameworks to promoting cross-border research and cooperation to avoid a new quantum divide, quantum diplomacy will be pivotal in shaping the future governance of quantum advancements. Two main areas were seen as crucial in this respect:

- a) The need for diplomatic action to make quantum technology available without hurdles, especially to ensure developing regions are not left behind by a new technology divide.
- b) Collaboration among scientists, industry, and policymakers to co-design frameworks for responsible innovation before quantum's dual-use dilemmas outpace governance.

Collaboration at international level can help countries grow their expertise, build a quantum-ready workforce, create opportunities (including for women and girls), equip youth for further quantum development, and foster quantum research locally.

International cooperation and early, inclusive governance are essential to steer quantum technologies toward widespread use and benefits. No single entity can build the quantum ecosystem alone – a global supply chain and standards are needed to ensure innovation is not siloed.

Together, diplomatic action to bridge the quantum divide, early adoption of technical standards, and global collaboration on workforce skills and ethics could form the backbone of a responsible, inclusive quantum future. Advancing Quantum for good will require balancing openness, security, and shared standards so that innovation is not siloed but genuinely benefits all.

22.5 The road to lasting impact

After a series of quantum-focused talks and panels, the wrap-up sessions attempted to answer the big question: *Quantum – hype, hope, or reality?* There was a unanimous sense that the field has unprecedented momentum in 2025 with record-breaking investments being made in quantum startups and national programs, steady progress in hardware (quantum bits are increasing, new error-correction milestones, etc.), and growing public awareness. Yet, experts cautioned that practical utility and scalability are the true measures of quantum technology's success – and by those measures, there is still a long road ahead. In other words, promise still outweighs proof in 2025.

The closing keynote offered a grounded but hopeful perspective: the quantum journey is a marathon, not a sprint. Drawing historical parallels, the keynote reminded the audience that the revolutionary impacts of classical computing did not happen overnight but were a result of decades of incremental advances in transistors, software, and engineering. Quantum technology may well follow a similar trajectory. Significant breakthroughs will require patience, sustained research, and coordinated investment across the "full stack" – from fundamental physics and materials science through to hardware engineering, software development, and algorithm design. Progress will also come in waves, with perhaps unforeseen Eureka moments along the way. "We don't yet know the killer application of quantum computing, and the only way to find it is to build these machines," the speaker noted, emphasizing the need for continued experimentation and innovation.

In conclusion, quantum technology's future is bright but must be approached with both enthusiasm and realism. The Quantum for Good track reinforced that tremendous possibilities are on the horizon if quantum science is harnessed for societal benefit but achieving this will demand global collaboration, policy guidance, and a steadfast focus on ethical, inclusive development. The next steps involve turning the insights from these discussions into action: launching collaborative projects, developing standards and governance frameworks, investing in education, and continuing to communicate the value of quantum innovation to the public.

and policymakers. With 2025 being a landmark year of quantum awareness and coordination, the stage is set for a new era where "quantum for good" can truly move from aspiration to reality.

22.6 Future directions

In this year of the International Year of Quantum Science and Technology, at a time when global attention is finally turning toward a field that has, for decades, remained largely invisible outside research labs, the question now is how to sustain that attention once the banners come down. Proactive international dialogue is essential to build trust and establish quantum standards for security, interoperability, and equitable global access, even before technology fully matures.

Four key announcements were made at this session signalling how ITU and partners plan to carry out the issues highlighted during the discussions:

- a) Extended Quantum for Good track in 2026: Building on the momentum from 2025, the Quantum for Good track will return in 2026 with an expanded program at the Quantum for Good Summit, featuring deeper engagement with global partners, a larger showcase, and new cross-sector challenges.
- b) UN-led demonstration of quantum technologies in action through a joint project between UNICC and ITU: A Quantum Key Distribution Network (QKDN) will be established between the ITU and UNICC facilities in Geneva, showcasing quantum-secure communication in real-world UN operations and ITU-T QKDN standards in action.
- c) Quantum World Tour: In partnership with The Quantum Insider, a global awareness-building series will kick off in September spotlighting innovation, startups, and talent across countries, bringing quantum dialogue and demos directly to communities worldwide
- d) Multi-lingual Introduction to Quantum course: A new, accessible *Introduction to Quantum* course will be launched in multiple languages to support broader global understanding of quantum science and technology. The course aims to democratize quantum literacy and reach underserved regions.

Acknowledgements

The organizers would like to acknowledge the speakers below for having kindly accepted to participate in the respective sessions of the International AI Standards Exchange during the AI for Good Summit 2025, for sharing their insights and for their excellent presentations.

Session	Speaker Name	Title	Organization
Centre Stage 11 July 2025			
The role of standards in shaping an AI-driven future	Seizo Onoe	Director of the Telecommunication Standardization Bureau (TSB)	International Telecommunication Union (ITU)
Role of International Standards in AI	Amandeep Singh Gill	Under-Secretary-General and Special Envoy for Digital and Emerging Technologies, Office for Digital and Emerging Technologies	United Nations
	Seizo Onoe	Director of the Telecommunication Standardization Bureau (TSB)	International Telecommunication Union (ITU)
	Sung Hwan Cho	President	International Organization for Standardization (ISO)
	Philippe Metzger	Secretary-General & CEO	International Electrotechnical Commission (IEC)
	Kathleen A. Kramer	President & CEO	IEEE
	Paul Gaskell	Deputy Director: Digital Trade, Internet Governance & Digital Standards	Department for Science, Innovation & Technology, UK
	Bilel Jamoussi	Deputy Director of TSB Chief, Study Groups & Policy Department (SPD)	International Telecommunication Union (ITU)
Transforming telecoms with AI and ML	John Omo	Secretary-General	African Telecommunications Union (ATU)
	Hatem Dowidar	Group CEO	e&
	Chih-Lin I	China Mobile Chief Scientist, Wireless Technologies	China Mobile Research Institute
	Pamela Snively	Chief Data & Trust Officer	TELUS
	Ebtesam Almazrouei	Executive Director of the Office of AI and Advanced Technology at the Department of Finance, CEO and Founder of AIE3	Chairperson of UN AI for Good Impact Initiative

Session	Speaker Name	Title	Organization
AI Standards: Building trust, ensuring innovation in networked world	Cameron F. Kerry	Ann R. and Andrew H. Tisch Distinguished Visiting Fellow	The Brookings Institution; Visiting Scholar, MIT Media Lab
Navigating tomorrow: Education, skills, and standards in an AI-driven workplace	Beena Ammanath	Author, Global Deloitte AI Institute Leader	Deloitte
Computing and AI: Endless frontiers and exploration	Wang Jian	Director, Founder of Alibaba Cloud	Zhejiang Lab
Donate your brainwaves for social impact	Rodrigo Hübner Mendes	CEO & Founder	Institute Rodrigo Mendes
	Olivier Oullier	Co-Founder & CEO, Inclusive Brains; Visiting Professor, Mohamed Bin Zayed University of Artificial Intelligence (MBZUAI); Chairman, AI Institute	Biotech Dental Group
Navigating deepfakes and securing multimedia authenticity in the age of generative AI	Alessandra Sala	Chair of AI and Multimedia Standards Collaboration, Sr. Director of Artificial Intelligence and Data Science	Shutterstock
AI for food systems	Seizo Onoe	Director of the Telecommunication Standardization Bureau (TSB)	International Telecommunication Union (ITU)
	Pieterneel Boogaard	Managing Director, Office of Technical Delivery	United Nations-International Fund for Agriculture Development (IFAD)
	Magan Naidoo	Chief Data Officer	United Nations World Food Programme (WFP)
	Dejan Jakovljevic	Chief Information Officer (CIO), Director of Digital FAO and Agro-Informatics Division	Food and Agriculture Organization of the United Nations (FAO)
	Sebastian Bosse	Head of the Interactive & Cognitive Systems Group	Fraunhofer Heinrich Hertz Institute (HHI)
	LJ Rich	Inventor/Musician and International Broadcaster	Perfect Pitch Productions

Session	Speaker Name	Title	Organization
ITU-WHO-WIPO Global Initiative on AI for Health – Future directions and launch of technical brief: Mapping the application of AI in Traditional Medicine	LJ Rich	Inventor/Musician and International Broadcaster	Perfect Pitch Productions
	Dalila Hamou	Director, External Relations Division	World Intellectual Property Organization (WIPO)
	Edward Kwakwa	Assistant Director General, Global Challenges and Partnerships Sector	World Intellectual Property Organization (WIPO)
	Seizo Onoe	Director of the Telecommunication Standardization Bureau (TSB)	International Telecommunication Union (ITU)
	Alain Labrique	Director for the Department of Digital Health and Innovation	World Health Organization (WHO)
AI and energy	Gitta Kutyniok	Bavarian AI Chair for Mathematical Foundations of Artificial Intelligence	Ludwig-Maximilians Universität München
	Qi Shuguang	Vice Deputy Engineer, China Telecommunication Technology Labs (systems)	China Academy of Information and Communications Technology (CAICT)
Future standards for frontier tech: quantum	Cierra Choucair	CEO, Universum Labs, Strategic Content Director	The Quantum Insider
	Qiang Zhang	Executive director, Jinan Institute of Quantum Technology, Professor	University of Science and Technology of China
Challenging the status quo of AI security	Xiaofang Yang	LLM Security Director	Ant Group
	Debora Comparin	Standardization Expert	Thales Digital Identity & Security
	Sounil Yu	Author and creator of the Cyber Defence Matrix and CTO	Knostic AI
Future of smart mobility	Helen Pan	General Manager	Baidu Apollo
	Vincent Vanhoucke	Distinguished Engineer	Waymo
	Francois E. Guichard	Focal Point, Intelligent Transport Systems and Automated Driving	United Nations Economic Commission for Europe (UNECE)

Session	Speaker Name	Title	Organization
Closing remarks	Doreen Bogdan-Martin	Secretary-General	International Telecommunication Union (ITU)

Thematic Workshops

AI for agriculture – Shaping standards for smart food systems	Daniel Young	Data Scientist and PhD Candidate	University of Texas, Austin
	Priya Samant	CEO and Co-Founder	Abris Inc.
	Pieterneel Boogaard	Managing Director, Office of Technical Delivery	United Nations-International Fund for Agriculture Development (IFAD)
	Brenda Mulele Gunde	Global Lead for ICT4D	United Nations-International Fund for Agriculture Development (IFAD)
	Henry van Burgsteden	Senior Innovation Officer (Office of Innovation)	Food and Agriculture Organization (FAO)
	Erik van Ingen	Digital Innovation Consultant	Food and Agriculture Organization (FAO)
	Olivier Francon	Associate Vice President, Neuro AI, AI Labs	Cognizant
	Dejan Jakovljevic	Chief Information Officer (CIO), Director of Digital FAO and Agro-Informatics Division	Food and Agriculture Organization (FAO)
	Mythili Menon	Project Officer Advisor	International Telecommunication Union (ITU)
	Sebastian Bosse	Head of the Interactive & Cognitive Systems Group	Fraunhofer Heinrich Hertz Institute (HHI)
	Risto Miikkulainen	VP of AI Research Cognizant AI Labs, Professor of Computer Science	University of Texas, Austin
	Christian Merz	Co-Lead ,FAIR Forward: Artificial Intelligence for All’ Division Economic and Social Development	Digitalisation
	Howard Lakougna	Senior Program Officer	Bill and Melinda Gates Foundation
	Alessandra Furtado	Director of Programs	International Potato Center (CIP)
	Kyriacos M. Koupparis	Head, Hunger Monitoring Unit	United Nations World Food Programme (WFP)

Session	Speaker Name	Title	Organization
	Magan Naidoo	Chief Data Officer	United Nations World Food Programme (WFP)
Enabling AI for health innovation and access	Rich Parker	Learning Expert & Emergencies Specialist	
	Peter Schiøler	Information Manager, WHO Health Emergencies Cairo	World Health Organization (WHO)
	Shyama Kuruvilla	Director, AI at the Global Centre for Traditional Medicine	World Health Organization (WHO)
	Claudia Seitz	Professor of Public Law, European Law, Public International Law and Life Sciences Law	UFL, Liechtenstein
	Komal Kalha	Deputy Director, Intellectual Property and Trade	IFPMA
	Nancy Pignataro	Associate External Relations Officer, External Relations Division	World Intellectual Property Organization (WIPO)
	Siddhartha Prakash	Head, Global Health, Global Challenges Division	World Intellectual Property Organization (WIPO)
	Tobias Schonwetter	Associate Professor, Department of Commercial Law	University of Cape Town
	Sneha Jain	IP litigator	Saikrishna and Associates, India
	Thierry Cordier Lassalle	Health Operations and Logistics Capacity Development, Health Emergency Programme	World Health Organization (WHO)
	Dalila Hamou	Director, External Relations Division	World Intellectual Property Organization (WIPO)
	Alain Labrique	Director for the Department of Digital Health and Innovation	World Health Organization (WHO)
	Yinzi Jin	Deputy Director, School of Public Health	Peking University
	Gopal Ramchurn	CEO	Reponsible AI UK
	Romita Ghosh	Founder and CEO	RevolutionAlze
	Mary-Anne ("Annie") Hartley	Professor Director	Laboratory for Intelligent Global Health & Humanitarian Response Technologies (LiGHT)

Session	Speaker Name	Title	Organization
	Zoltán Turbék	Deputy Permanent Representative	Permanent Mission of Hungary to the United Nations Office and other international organizations in Geneva
	Simão Campos	Counsellor, AI for Health standards	International Telecommunication Union (ITU)
	André Anjos	Group Leader in Medical AI and Scientific Researcher	Idiap Research Institute, Switzerland
	Nina Linder	Co-Chair of Topic Group point-of-care diagnostics	University of Helsinki
	Sameer Pujari	Lead AI for health	World Health Organization (WHO)
	Bilel Jamoussi	Deputy Director of TSB Chief, Study Groups & Policy Department (SPD)	International Telecommunication Union (ITU)
Human-centered AI for disaster management: Empowering communities through standards	Ryuki Hyodo	Chief Science Officer	SpaceData Inc.
	Jumpei Takami	Associate Expert in Remote Sensing for Disaster Management, UN-SPIDER Program	United Nations Office for Outer Space Affairs (UNOOSA)
	Hwirin Kim	Chief of Hydrological and Water Resources Services Section	World Meteorological Organization (WMO)
	Leonardo Milano	Data Science Team Lead	United Nations Office for the Coordination of Humanitarian Affairs (UNOCHA)
	Claudio Rossi	Program manager & Senior researcher	LINKS Foundation
	Taeyoon Kim	Programme Officer	United Nations Framework Convention on Climate Change (UNFCCC) Secretariat
	Alen Berta	PMP, Executive Space Consultant and the Head of CGI AI and Advance Analytics Competence Area within CGI Germany operations	CGI Deutschland B.V. & Co. KG

Session	Speaker Name	Title	Organization
	Fumiko Nohara	Senior Expert for Emergencies and Postal Resilience, Development Cooperation Directorate (DCDEV)	The Universal Postal Union (UPU)
	Giriraj Amarnath	Principal Researcher – Disaster Risk Management and Climate Resilience & CGIAR Interim Deputy Director for Climate Action	International Water Management Institute (IWMI), Sri Lanka
	Katharina Weitz	Project Manager	Fraunhofer Heinrich Hertz Institute (HHI)
	Kevin White	Senior Director, AI for Good Lab	Microsoft
	Pierre-Philippe Mathieu	ESA Implementation Manager, Civil Security from Space (CSS)	European Space Agency (ESA)
	Markus Reichstein	Director & Professor	Max Planck Institute for Biogeochemistry
	Monique Kuglitsch	Innovation Manager	Fraunhofer Heinrich Hertz Institute (HHI)
	Muralee Thummarukudy	Director of the Coordination Office of the G20 Initiative on Land	United Nations Convention to Combat Desertification (UNCCD)
	Carlos Uribe	Programme Specialist Disaster Risk Reduction Unit - Natural Sciences Sector	United Nations Educational, Scientific and Cultural Organization (UNESCO)
Navigating the intersect of AI, environment and energy for a sustainable future	Eugênio Vargas Garcia	Ambassador and Director, Department of Science, Technology, Innovation and Intellectual Property	Ministry of Foreign Affairs of Brazil
	Jonathan Turnbull	Program Manager - gTech Sustainability	Google
	Ivana Drobnjak	Director of the Undergraduate Programme in Computer Science	University College London (UCL)
	Dietram Oppelt	Chair	UNFCCC Technology Executive Committee (TEC)
	John Omo	Secretary-General	African Telecommunications Union (ATU)
	Sara Ballan	Senior Digital Development Specialist, Co-lead Green Digital Business Line	The World Bank
	Carlos Souza Jr.	Associate researcher	Amazon Institute of People and the Environment

Session	Speaker Name	Title	Organization
	Tim Smolcic	Director IT Sustainability Strategy	HSBC
	Dominique Würges	Chair, ITU-T Study Group 5 , International Telecommunication Union (ITU), Director of Institutional Relations for Standardisation	Orange
	Juliette Fropier	Project Manager, AI & Ecological Transition	Ministries of Ecology, Energy, and Territories, France
	Roy Schwartz	Associate professor	The Hebrew University of Jerusalem
	Li Cui	Chief Development Officer and Chief of Staff to CEO	ZTE Corporation
	Priyadarshini Panda	Assistant Professor of Electrical & Computer Engineering	Yale University
	Rosendo Manas	Co-founder	Resilio
	Laure de la Raudière	President	Autorité de Régulation des Communications électroniques et des Postes (ARCEP), France
	Josh Parker	Senior Director, Corporate sustainability	NVIDIA
	Francesca Dominici	Clarence James Gamble Professor of Biostatistics, Population, and Data Science, Biostatistics, Harvard T.H. Chan School of Public Health	Harvard University
	Björn Ommer	Head of Computer Vision & Learning Group	Ludwig-Maximilians Universität München
	Amine Saboni	MLOps Engineer	Pruna AI
	Yolanda Martinez	Practice Manager for Digital Development in Latin American and the Caribbean	The World Bank
	Philippe Tuzzolino	SG5 expert and Vice President Environment	Orange
	Lina María Duque Del Vecchio	Executive Director and Commissioner	Commission for Communications Regulation of Colombia (CRC)
	Tomas Lamanauskas	Deputy Secretary-General	International Telecommunication Union (ITU)

Session	Speaker Name	Title	Organization
	Antonio De Domenico	Principal Engineer	Huawei Technologies Co., Ltd.
	Gitta Kutyniok	Bavarian AI Chair for Mathematical Foundations of Artificial Intelligence	Ludwig-Maximilians Universität München
AI readiness – Towards a standardized readiness framework	Hoda Baraka	Advisor to Minister of ICT for Technology Talents Development, National AI Lead and Acting Director	Egyptian Center for Responsible AI
	Anna Abramova	Director	MGIMO AI Centre
	Isabella de Michelis di Slonghello	CEO and Founder	Ernieapp ltd
	Hong-Chuan Yang	Professor, Department of Electrical and Computer Engineering	University of Victoria
	Lindsey Moore	CEO and Founder	DevelopMetrics
	Mohammed Alawad	General Manager of the General Administration of Studies	Saudi Data and AI Authority (SDAIA)
	Rehab Alarfaj	General Manager of Strategic Partnerships and Indices	Saudi Data and Artificial Intelligence Authority (SDAIA)
	Marcelo Mendoza	Associate professor, Computer Science Department	Pontificia Universidad Católica de Chile
	Daniel Rezende	R&D Program Manager	Lenovo
	Osmar Bambini	Climate innovation ecopreneur, Co-founder and CIO	umgrauemeio (1.5°C)
	Ahmed Biyabani	Associate teaching professor	Carnegie Mellon University
	Francesco Flammini	Professor of Trustworthy Autonomous Systems	The University of Applied Sciences and Arts of Southern Switzerland (SUPSI)
	Kai Wei	Vice Chair of SC on AI; Director, Artificial Intelligence Research Institute	China Academy of Information and Communications Technology (CAICT)
	Lilibeth Acosta	Deputy Director, Climate Action and Inclusive Development Department	Global Green Growth Institute (GGGI)
	Chenxi Qiu	Deputy Director, Center for Global Digital Governance (CGDG)	China Academy of Information and Communications Technology (CAICT)
	Yu Xiaohui	President	China Academy of Information and

Session	Speaker Name	Title	Organization
			Communications Technology (CAICT)
	Asrat Mulatu Beyene	Associate Professor	Addis Ababa Science and Technology University
	Jiaying Meng	Junior Programme Officer	International Telecommunication Union (ITU)
	Lina Bariah	Senior Researcher	Khalifa University of Science and Technology
	James Agajo	Professor	Federal University of Technology Minna Nigeria
	Thomas Basikolo	Programme Coordinator	International Telecommunication Union (ITU)
	Vishnu Ram OV	Independent Research Consultant	International Telecommunication Union (ITU)
	Abdukodir Khakimov	Chief IT Architect	RUDN University IONAT Tajikistan
Trustworthy AI testing and validation	Bilel Jamoussi	Deputy Director of TSB Chief, Study Groups & Policy Department (SPD)	International Telecommunication Union (ITU)
	Dawn Song	Professor, Computer Science @ UC Berkeley and Director	Berkeley RDI (Berkeley center for responsible decentralized intelligence)
	Hiromu Kitamura	Principal Expert for Technical Management	Japan AI Safety Institute
	Robert Trager	Co-Director, Oxford Martin AI Governance Initiative	University of Oxford
	Lam Kwok Yan	Professor of Computer Science, College of Computing and Data Science	Nanyang Technological University Singapore
	Sray Agarwal	Head of Responsible AI (EMEA & APAC)	Infosys
	Sebastian Hallensleben	Chair	CEN-CENELEC JTC 21
	Franziska Weindauer	CEO	TÜV AI.Lab
	Evan Miyazono	CEO	Atlas Computing

Session	Speaker Name	Title	Organization
	Karine Perset	Acting Head of the OECD Division on AI and Emerging Digital Technologies	Organization for Economic Cooperation and Development (OECD)
	Richard Kerkdijk	Senior Consultant Cyber Security Technologies	TNO
	Puck de Haan	Cyber Security & AI Researcher	TNO
	Tammy Masterson	Head of Best Practices	UK AI Safety Institute
	Vijay Mauree	Programme Coordinator, TSB	International Telecommunication Union (ITU)
	Jaime Fernández Fisac	Assistant Professor of Electrical and Computer Engineering	Princeton University
	Mark Nitzberg	Interim Executive Director and co-founder	International Association for Safe & Ethical AI (IASAI)
	Chris Meserole	Executive Director	Frontier Model Forum
Open dialogue for trustworthy AI testing	Seizo Onoe	Director of the Telecommunication Standardization Bureau (TSB)	International Telecommunication Union (ITU)
	Bilel Jamoussi	Deputy Director of TSB Chief, Study Groups & Policy Department (SPD)	International Telecommunication Union (ITU)
	Robert Trager	Co-Director, Oxford Martin AI Governance Initiative	University of Oxford
	Vijay Mauree	Programme Coordinator, TSB	International Telecommunication Union (ITU)
	Renan Araujo	Research Manager, International AI Strategy	Institute for AI Policy and Strategy (IAPS)
	Sam Daws	Senior Advisor, Oxford Martin AI Governance Initiative	University of Oxford
	Sumaya Adan	AI Governance Researcher, Oxford Martin AI Governance Initiative	University of Oxford
AI and Machine Learning in communication networks 9-Jul	Stephen Kolesh	Machine Learning Engineer	Sportserve
	Jeanine Vos	Head of SDG Accelerator	GSMA
	Murat Unlusan	Associate Director, Access Network Architecture	Turkcell

Session	Speaker Name	Title	Organization
	Hoda Baraka	Advisor to Minister of ICT for Technology Talents Development, National AI Lead and Acting Director	Egyptian Center for Responsible AI
	David Cox	VP for AI models	IBM Research
	Daniel Becking	Research associate and Ph.D. student	Fraunhofer HHI and TU Berlin
	Zhiyuan Jiang	Professor	Shanghai University
	Nada Golmie	NIST Fellow	National Institute of Standards and Technology (NIST)
	Lina Bariah	Senior Researcher	Khalifa University of Science and Technology
	Nicola Piovesan	Senior Researcher	Huawei
	Antonio De Domenico	Principal Engineer	Huawei Technologies Co., Ltd.
	Ahmed Alkhateeb	Assistant Professor	Arizona State University
	Wei Meng	Director of standard and open source planning	ZTE Corporation
	Henning Sanneck	Research Manager	EU SNS-JU 6G-MIRAI-HARMONY project
	Vishnu Ram OV	Independent Research Consultant	International Telecommunication Union (ITU)
	Chih-Lin I	China Mobile Chief Scientist, Wireless Technologies	China Mobile Research Institute
	Thomas Basikolo	Programme Coordinator	International Telecommunication Union (ITU)
	Bilel Jamoussi	Deputy Director of TSB Chief, Study Groups & Policy Department (SPD)	International Telecommunication Union (ITU)
	Celina Lee	Co-Founder & CEO	Zindi Africa
	Wafae Bakkali	Staff Generative AI Specialist, Blackbelt	Google
AI and Machine Learning in communication networks 10-Jul	Shota Ono	Research Associate	The University of Tokyo
	Sebastian Cammerer	Senior Research Scientist	NVIDIA
	Riccardo Trivisonno	Head of Network Architecture	Huawei Technologies Co., Ltd.

Session	Speaker Name	Title	Organization
	Marco Carugi	Senior Consultant, Advanced ICTs and Standardization	Huawei Technologies Co., Ltd.
	Buse Bilgin	Next Generation R&D Engineer	Turkcell
	Jose M. Alcaraz Calero	Professor	Aston University, UK
	Qi Wang	Professor	University of Leicester
	Nada Golmie	NIST Fellow	National Institute of Standards and Technology (NIST)
	Nicola Piovesan	Senior Researcher	Huawei
	Antonio De Domenico	Principal Engineer	Huawei Technologies Co., Ltd.
	Gyu Myoung Lee	Professor	Liverpool John Moores University (LJMU)
	Liya Yuan	Open Source & Standardization Engineer	ZTE Corporation
	Francesc Wilhelmi	Professor	Universitat Pompeu Fabra (UPF)
	Vishnu Ram OV	Independent Research Consultant	International Telecommunication Union (ITU)
	Chih-Lin I	China Mobile Chief Scientist, Wireless Technologies	China Mobile Research Institute
From policy to practice: Building trust in multimedia authenticity through international standards	Alessandra Sala	Global President of Women in AI; Chair of AI and Multimedia Standards Collaboration, Sr. Director of Artificial Intelligence and Data Science	Shutterstock
	Touradj Ebrahimi	Professor, Ecole Polytechnique Fédérale de Lausanne, Convenor of JPEG Group, and Vice Chair	AI and Multimedia Authenticity Standards Collaboration
	Leonard Rosenthol	Senior Principal Architect	Adobe
	Jacobo Castellanos	Coordinator Technology, Threats, and Opportunities	WITNESS
	Carol Buttle	International Advisor to Governments on Technology and Critical National Infrastructure	
	Cindy Parokkil	AI Policy Lead & Programme Manager	ISO Central Secretariat

Session	Speaker Name	Title	Organization
Future Networked Car symposium 2025	Seizo Onoe	Director of the Telecommunication Standardization Bureau (TSB)	International Telecommunication Union (ITU)
	Christian Thiele	Senior Director, Global Ground Vehicle Standards	SAE International
	Arnd Bätzner	Founder	BAETZNER METROPOLITAN
	Harald Barth	Product Marketing Manager, Driving Assistance	Valeo Group
	Edward Wilford	Senior Research Director, Automotive	Omdia
	Alexey Rybakov	VP, Automotive Products	Netradyne
	Simone Fabris	VP of Product & Delivery	Wayve
	Paula Palade	AI Ethics Senior Technical Specialist	Jaguar Land Rover (JLR)
	T. Russell Shields	Chair	CITS
	Helen Pan	General Manager	Baidu Apollo
	Katherine Evans	Chief Representative of the IEEE to the UNECE	IEEE SA
	Jose Ignacio (Nacho) Rexach	Chief Commercial Officer (CCO) for Europe and Latin America	EHang (NASDAQ: EH)
	William Gouse	Director, Federal Program Development	SAE International
	Bilel Jamoussi	Deputy Director of TSB Chief, Study Groups & Policy Department (SPD)	International Telecommunication Union (ITU)
	Alexandre Corjon	Senior VP & Technical Fellow	Sonatus
	Roger C. Lanctot	President	Mobile Satellite Users Association
	Hongki Cha	Special Fellow	Electronics and Telecommunications Research Institute (ETRI)
	Stefano Polidori	Counsellor, ITU-T Study Group 21	International Telecommunication Union (ITU)
	Dmitry Mariyasin	Deputy Executive Secretary	United Nations Economic Commission for Europe (UNECE)
	Vincent Vanhoucke	Distinguished Engineer	Waymo
	Francois E. Guichard	Focal Point, Intelligent Transport Systems and Automated Driving	United Nations Economic Commission for Europe (UNECE)

Session	Speaker Name	Title	Organization
	Jean Todt	United Nations Special Envoy for Road Safety	United Nations
Empowering innovative and intelligent solutions at the edge	David Nahmani	Co-founder and CRO	MountAln
	Jonathan Russo	Program Manager	femtoAI
	Saharsh Singhania	AVP Product Marketing	Ambient Scientific
	Moritz Joseph	Co-founder and CEO	RooflineAI GmbH
	Shahnawaz Ahmed	Deep Learning Researcher	EmbedI
	Eiman Kanjo	Professor, TinyML and Pervasive Computing and Head of the Smart Sensing Lab	Nottingham Trent University
Empowering innovative and intelligent solutions at the edge	Avijit Sinha	Senior Vice President, Strategy and Global Business Development	Wind River
	Ning Chen	Chairman and CEO	Intellifusion
	Marco Zennaro	Research Scientist	Abdus Salam International Centre for Theoretical Physics
	Bilel Jamoussi	Deputy Director of TSB Chief, Study Groups & Policy Department (SPD)	International Telecommunication Union (ITU)
	Nadim Maamari	Head of Edge AI and Vision System	CSEM
	Marie Didier	Founder	MATIS
Challenging the status quo of AI security	Abbie Barbir	SG/17 Q10/17 Co-Rapporteur	International Telecommunication Union (ITU)
	Alan Chan	Research Fellow	Centre for the Governance of AI
	Tobin South	AI Research Fellow at Stanford University Head of AI Agents	WorkOS
	Xiao Rong	Vice President and General Manager, AI Technology Platform	Shenzhen Intellifusion Technologies Co., Ltd.
	Evan Miyazono	CEO	Atlas Computing
	Xiaofang Yang	LLM Security Director	Ant Group
	Lisa Bechtold	Head of AI Governance	Zurich Insurance Group

Session	Speaker Name	Title	Organization
	Jabu Mtsweni	Chief Researcher and Centre Manager for the Information and Cybersecurity Research Centre	Council for Scientific and Industrial Research (CSIR) South Africa
	Kishor Narang	Mentor, Principal Design Strategist & Architect	Narnix Technolabs Pvt. Ltd
	Debora Comparin	Standardization Expert	Thales Digital Identity & Security
	Naying Hu	Senior Business Executive of AI Institute	China Academy of Information and Communications Technology (CAICT)
	Arnaud Taddei	Global Security Strategist	Broadcom
	Sounil Yu	Author and creator of the Cyber Defence Matrix and CTO	Knostic AI
	Kai Wei	Vice Chair of SC on AI; Director, Artificial Intelligence Research Institute	China Academy of Information and Communications Technology (CAICT)
	Babak Hodjat	Chief Technology Officer AI	Cognizant
	Kay Firth-Butterfield	CEO	Good Tech Advisory
	Francesca Bosco	Chief Strategy Officer	CyberPeace Institute
Women leaders in AI & standards	H.E. Ms. Neema Lugangira	Women Political Leaders' Secretary General Member of Parliament	Tanzania (United Republic of)
	Doreen Bogdan-Martin	Secretary-General	International Telecommunication Union (ITU)
	Kathleen A. Kramer	President & CEO	IEEE
	Alessandra Sala	Global President of Women in AI; Chair of AI and Multimedia Standards Collaboration, Sr. Director of Artificial Intelligence and Data Science	Shutterstock
	Roser Almenar	ITU Secretary-General Youth Advisory Board member, PhD Candidate	AI & Space Law at the University of Valencia
	Salma Abbasi	Founder, Chairperson and CEO	eWorldwide Group
	Rim Belhassine-Cherif	Chair, Network of Women in ITU-T (NoW in ITU-T), Chief Innovation and Strategy Officer	Tunisie Télécom, Tunisia
	Paola Cecchi Dimeglio	Chair, Executive Leadership Research Initiative for Women and Minority (ELRIWMA); Professor	Harvard University
	Celina Lee	Co-Founder & CEO	Zindi Africa

Session	Speaker Name	Title	Organization
	Charlyne Restivo	Programme Coordinator	International Telecommunication Union (ITU)
Quantum for Good: Industry leadership, innovation and real-world impact	Deepa Shinde	Lead Pharma & Life Sciences Architect & Quantum Computing and Quantum Safe Ambassador	Microsoft
	Zina Jarrahi Cinker	Director General at MATTER Chief Creator	XPANSE
	Mekena McGrew	Quantum musician	
	Claudius Riek	Managing Director	Zurich Instruments Germany
	Wiktor Mazin	Quantum fractal artist	
	Debora Comparin	Standardization Expert	Thales Digital Identity & Security
	Graham Alabaster	Director Geneva Office	UN-Habitat
	Anne Dames	Distinguished Engineer, IBM Infrastructure & IBM Z Cryptographic Technology Development	IBM
	Gregoire Ribordy	CEO	ID Quantique (IDQ)
	Katsuyuki Hanai	Chair of Quantum Cryptography and Quantum Communication subcommittee, Quantum STRategic industry Alliance for Revolution (Q-STAR) Business Unit Manager, ICT Solutions Division	Toshiba Digital Solutions Corp
	Richard Hall-Wilton	Director at Center for Sensors & Devices	Fondazione Bruno Kessler (FBK), Italy
	Mathias Soeken	Principal Quantum Software Architect	Microsoft
	Mio Murao	Professor, Department of Physics	University of Tokyo
	Yaseera Ismail	Senior Lecturer, Department of Physics	Stellenbosch University, South Africa
	Joanna Doummar	Associate Professor of Groundwater Hydrology	American University of Beirut in Lebanon
	Catherine Lefebvre	Senior Advisor at Geneva Science and Diplomacy Anticipator (GESDA)	Open Quantum Institute (OQI)
	Elif Kiesow	Executive Director, Ethicqual and Ethics Subcommittee Lead	Quantum Economic Development Consortium (QED-C)

Session	Speaker Name	Title	Organization
Quantum technology and diplomacy: Why should we care?	Will Zeng	Partner, Quantonation and President	Unitary Foundation
	Leandro Aolita	Acting Chief Researcher, Quantum Research Center	Technology Innovation Institute (TII)
	René Reimann	Senior Director, Quantum Sensing	Technology Innovation Institute (TII)
	Ulrich Mans	Strategic Partnerships Lead	Quantum Delta, Netherlands
	Mira Wolf-Bauwens	Head of Initiatives Development, Geneva Science and Diplomacy Anticipator (GESDA)	Open Quantum Institute (OQI)
	Thierry Botter	Executive Director	European Quantum Industry Consortium (QuIC)
	Najwa Aaraj	CEO	Technology Innovation Institute (TII)
	Sameer Chauhan	Director	United Nations International Computing Centre (UNICC)
	Seizo Onoe	Director of the Telecommunication Standardization Bureau (TSB)	International Telecommunication Union (ITU)
	Carlos Ruiz-Garvia	Team Lead, Adaptation Committee Unit	United Nations Framework Convention on Climate Change (UNFCCC) Secretariat
	Carlos Abellán	cofounder and CEO	Quside
	Cierra Choucair	CEO, Universum Labs, Strategic Content Director	The Quantum Insider
	Eleni Diamanti	Research Director, Sorbonne Université and Co-founder	WelinQ
	H.E. Mr. Omar Zniber	Ambassador	Permanent Representative of the Kingdom of Morocco to the United Nations Office and other International Organizations in Geneva

Session	Speaker Name	Title	Organization
Quantum for all: Innovation, access & impact	Maricela Muñoz	Director of External Affairs	Geneva Science and Diplomacy Anticipator (GESDA)
	Justine Lacey	Director, Responsible Innovation	Commonwealth Scientific & Industrial Research Organisation (CSIRO), Australia
	Clare Shelley-Egan	Associate Professor, Ethics of Quantum Technologies	TU Delft
	Diederick Croese	Director	Centre for Quantum & Society, Quantum Delta, Netherlands
Building tomorrow's quantum workforce: Voices of future leaders in quantum	Mira Wolf-Bauwens	Head of Initiatives Development, Geneva Science and Diplomacy Anticipator (GESDA)	Open Quantum Institute (OQI)
	Caden Kacmarynski	Quantum Coalition Chief AI Architect	Moreland Connect
	Rui Rui Xie	Junior Project Officer, International Telecommunication Union (ITU)	Graduate Student, University College London
	Thomas Verrill	Quantum Coalition	Princeton University
	Dmitrii Khitrin	Quantum Coalition	Quantum Coalition
	Ben McDonough	Alumnus Advisor	University of Colorado Boulder
Quantum for Good: Shaping the future of quantum – What happens next? Future standards for frontier tech: quantum	Leandro Aolita	Acting Chief Researcher, Quantum Research Center	Technology Innovation Institute (TII)
	Cierra Choucair	CEO, Universum Labs, Strategic Content Director	The Quantum Insider
	Qiang Zhang	Executive director, Jinan Institute of Quantum Technology, Professor	University of Science and Technology of China
Brain-Computer interface: Key technical standards and diverse application scenarios	Yihan Wang	Marketing Manager, European Division	iFLYTEK
	Hervé Chneiweiss	Director, Neuroscience-Paris Seine-IBPS Laboratory ICNRS INSERM	Sorbonne University

Session	Speaker Name	Title	Organization
	Liyan Liang	Researcher	China Academy of Information and Communications Technology (CAICT)
	Dimiter Prodanov	Associate Professor	Bulgarian Academy of Sciences
	Xinyi Gu	Assistant Researcher	Beijing Institute of Technology
	Gerwin Schalk	Professor	Fudan University, Shanghai
	Ryota Kanai	Founder	Araya, Inc.
	Zhijin Qin	Associate Professor, Department of Electronic Engineering	Tsinghua University
	Yihua Ning	Founder and CEO	SceneRay
	Xie Songyun	Full Professor	Northwestern Polytechnical University
	Fabio Babiloni	Full professor of Physiology	University of Rome “Sapienza”
	Li Wenyu	Director Intellectual Property and Innovation Development Center	China Academy of Information and Communications Technology (CAICT)
	Saurabh Bhaskar Shaw	Adjunct Research Professor	Western University
	Frederic Werner	Chief, Strategy and Operations, AI for Good, TSB	International Telecommunication Union (ITU)
	H.E. Mr. Shan Zhongde	Vice Minister	Ministry of Industry and Information Technology, People’s Republic of China
	Seizo Onoe	Director of the Telecommunication Standardization Bureau (TSB)	International Telecommunication Union (ITU)
	Yu Xiaohui	President	China Academy of Information and Communications Technology (CAICT)
Building the technical foundations for embodied intelligence in connected ICT environments	Patricia Shaw	CEO	Beyond Reach Consulting Limited

Session	Speaker Name	Title	Organization
	Abhishek Gupta	CEO	Open Droids Robotics
	Zhang Min	European Regional Director	Unitree Robotics
	Kai Wei	Vice Chair of SC on AI; Director, Artificial Intelligence Research Institute	China Academy of Information and Communications Technology (CAICT)
	Selma Šabanović	Full Professor of Informatics and Cognitive Science	Indiana University Bloomington
	Guillem Martínez Roura	AI and Robotics Programme Officer	International Telecommunication Union (ITU)
	Noah Luo	Chair of ITU-T SG 21	International Telecommunication Union (ITU)
AI & virtual worlds: Building the cities and governments of tomorrow	Jimena Luna	Global AI Policy Lead	Lenovo
	Silvia Grandi	Director	Department for Digital, Connectivity and New Technologies, Ministry of Enterprises and Made in Italy
	Idah Pswarayi-Riddihough	Global Director and Director of Strategy and Operations at the Digital Transformation Vice-Presidency	World Bank Group
	H.E. Mr. William Kabogo Gitau	Cabinet Secretary	Ministry of Information, Communications and The Digital Economy (MICDE), Kenya
	Lise Nicoud	Deputy Mayor	City of Évian, France
	Luis Santiago	Innovation Adviser	Government Office for Information Technologies and Communication (OGTIC)
	H.E. Mr Jerry William Silaa	Minister	Ministry of Information, Communication and Information Technology, Tanzania (United Republic of)

Session	Speaker Name	Title	Organization
	Jeong-Kee Kim	Secretary General	World Smart Sustainable Cities Organization (WeGO)
	Muath AlRumayh	GM of International Affairs (CST) and Co-chair of the Joint Coordination Activity on Metaverse	International Telecommunication Union (ITU)
	Kamelia Kemileva	Co-Director, Programmes and Administration	Global Cities Hub
	Paola Cecchi Dimeglio	Chair, Executive Leadership Research Initiative for Women and Minority (ELRIWMA); Professor	Harvard University
	Bertrand Levy	Vice President Partnerships	The Sandbox
	Pilar Orero	Professor	Universitat Autònoma de Barcelona, Spain
	Martin Brynskov	Founding Director & Standardisation Lead	Open & Agile Smart Cities (OASC)
	Michelle Gyles-McDonnough	Executive Director	United Nations Institute for Training and Research (UNITAR)
	Raquel Brizída Castro	Vice-Chairwoman	Autoridade Nacional de Comunicações (ANACOM), Portugal
	Claudia Ximena Bustamante	Executive Director	Comisión de Regulación de Comunicaciones (CRC), Colombia
	Maria Galindo Garcia-Delgado	Secretary of Digital Policies	Government of Catalonia
	H.E. Mrs. Rina Yessenia Lozano Gallegos	Ambassador Extraordinary and Plenipotentiary	Permanent Mission of the Republic of El Salvador to the United Nations Office and other international organizations in Geneva
	Sameer Chauhan	Director	United Nations International Computing Centre (UNICC)
	Martin Wählisch	Associate Professor, Transformative Technologies, Innovation and Global Affairs	University of Birmingham, United Kingdom

Session	Speaker Name	Title	Organization
AI and Multimedia Authenticity Standards Workshop	Seizo Onoe	Director of the Telecommunication Standardization Bureau (TSB)	International Telecommunication Union (ITU)
	Okan Geray	Chair	Steering Committee of the Global Initiative on Virtual Worlds and AI
	Cristina Bueti	Counsellor	International Telecommunication Union (ITU)
	Hyoung Jun Kim	Chair, ITU-T Study Group 20 – Internet of Things, digital twins and smart sustainable cities and communities	International Telecommunication Union (ITU)
	Ambrose Ruyooka	Head of Research and Development	Ministry of ICT & National Guidance, Uganda
	Tomaz Levak	Co-founder	Umanitek
	Carol Buttle	International Advisor to Governments on Technology and Critical National Infrastructure	
	Jin Peng	Deputy General Manager of Technology Strategy and Execution Division	Ant Group
	Philippe Rixhon	Chair of the Management Board	Valunode OÜ, Tallinn, Estonia
	Farzaneh Badiei	Founder	Digital Medusa
	Jo Levy	Chair, Partner	Alliance for Responsible Data Collection, The Norton Law Firm
	Artemis Seaford	Head of AI Safety	ElevenLabs
	Cindy Parokkil	AI Policy Lead & Programme Manager	ISO Central Secretariat
	Sam Gregory	Human rights technologist, TED AI and deepfakes speaker, Executive Director	WITNESS
	Maurice Turner	Technical Policy Lead	TikTok

Session	Speaker Name	Title	Organization
	Mike Mullane	Deputy Director of Communications	International Electrotechnical Commission (IEC)
	Touradj Ebrahimi	Professor, Ecole Polytechnique Fédérale de Lausanne, Convenor of JPEG Group, and Vice Chair	EPFL
	Leonard Rosenthol	Senior Principal Architect	Adobe
	Alessandra Sala	Global President of Women in AI; Chair of AI and Multimedia Standards Collaboration, Sr. Director of Artificial Intelligence and Data Science	Shutterstock
	Thomas Wiegand	Executive Director	Fraunhofer Heinrich Hertz Institute (HHI)
	Amir Banifatemi	Chief Responsible AI Officer	Cognizant
	Vijay Mauree	Programme Coordinator, TSB	International Telecommunication Union (ITU)