

Rome Declaration for an Unarmed and Disarming Peace 2026, Mandate

August 13, 2026 • vajiramandravi.com



The **Rome Declaration for an Unarmed and Disarming Peace** was signed in July 2026 by Nobel laureates, AI scientists, religious leaders and other global figures. It warns about the growing risks of delegating moral and ethical decisions to Artificial Intelligence. Its central concern is nuclear command and control, where autonomous AI could reduce human decision making time and create catastrophic consequences during crises.

What is Rome Declaration for an Unarmed and Disarming Peace?

The Rome Declaration for an Unarmed and Disarming Peace is a moral and strategic appeal for human control over AI and renewed nuclear disarmament. It was adopted by participants of the Global Nobel Laureates Assembly. It is inspired by Pope Leo XIV's encyclical *Magnifica Humanitas* and seeks international safeguards against autonomous nuclear decisions and AI enabled warfare.

Rome Declaration Need and Purpose

The Rome Declaration for an Unarmed and Disarming Peace responds to faster military automation, weakening arms control arrangements and the possibility that AI errors could trigger irreversible decisions.

- **Shorter decision windows:** AI battle management systems can reduce military response time from hours to seconds, removing the human hesitation that helped prevent escalation during crises such as the Cuban Missile Crisis of 1962.
- **AI failure risks:** Deep neural networks can face hallucinations, data poisoning, adversarial attacks and faulty information, creating risks of false identification during high stakes military confrontations.
- **Weakening arms control:** The weakening of arrangements such as Open Skies and New START increases strategic uncertainty and may encourage states to automate defence systems to match perceived technological advantages.
- **Nuclear escalation:** In volatile situations involving nuclear programmes, cyber intrusions or physical attacks could be misread by automated systems and produce disproportionate military responses before human communication occurs.

Rome Declaration 6 Principles

The six principles of Rome Declaration for an Unarmed and Disarming Peace establish a broad framework linking AI safety, responsible governance, ethical leadership, military restraint and complete nuclear disarmament.

1. **Disarming the Next Arms Race:** States should prevent simultaneous arms races involving AI and nuclear capabilities and reduce incentives for automated military competition.
2. **Responsible Development:** AI development should follow transparent ethical standards, international principles, human rights, safety requirements and meaningful accountability.
3. **Responsible Use of AI:** AI should support human welfare rather than independently control nuclear launches, lethal force, military targeting, or autonomous warfare.
4. **Responsible Governance:** Governments should establish international oversight, regulatory mechanisms, accountability systems and verification arrangements for high risk AI applications.
5. **Responsible Leadership:** Governments, technology companies, scientists and military institutions should place human dignity, global peace and ethical responsibility above technological competition.
6. **Nuclear Disarmament:** The Declaration renews the objective of completely, irreversibly and verifiably eliminating nuclear weapons through sustained international negotiations.

Rome Declaration 5 Mandates

Five operational mandates translate the ethical concerns of Rome Declaration for an Unarmed and Disarming Peace into specific safeguards for AI governance, nuclear security, public knowledge and disarmament.

1. **Meaningful Human Control:** AI systems must never receive final authority over nuclear weapon launches, armed conflict initiation, or lethal force decisions, keeping responsibility with human decision makers.
2. **Digital Commons Model:** Greater access to relevant data, research and safety assessments should help independent experts examine systemic AI risks and unintended consequences.

3. **Responsible AI Development:** Developers should publish ethical frameworks and ensure systems remain explainable, auditable, controllable and capable of human override or shutdown.
4. **Internal Arsenal Vulnerability Audits:** Nuclear armed states should conduct rigorous technical reviews of command and control networks to identify AI driven cyber interference, data poisoning, hallucinations and other vulnerabilities.
5. **Time Bound Verifiable Disarmament:** States should resume good faith negotiations aimed at permanent, irreversible and verifiable elimination of nuclear arsenals.

Convergent Perils in AI

A convergent peril occurs when separate technological and governance risks interact and amplify one another into a larger systemic or civilizational threat.

- **Multiple risks interacting:** Algorithmic bias, hyper automation, cyber vulnerabilities, weak accountability and autonomous decision making can reinforce one another instead of remaining isolated problems.
- **Moral agency erosion:** Outsourcing ethical choices to machines may weaken human responsibility because AI cannot experience empathy, suffering, moral hesitation, or personal accountability.
- **Value alignment problem:** AI may optimise mathematical objectives efficiently, while human justice also considers rights, fairness, dignity, duties and individual circumstances.
- **Black box accountability:** Advanced AI systems can make decisions that are difficult to explain, creating uncertainty about whether developers, operators, governments, manufacturers, or users should bear responsibility.
- **Contextual blindness:** Human decision makers can consider compassion, intent, exceptional circumstances and social context, while AI systems may apply rules without understanding these human nuances.

AI Ethics in India

India has developed a principle based approach that seeks to combine technological innovation with inclusion, safety, transparency, privacy, fairness and accountability.

- **National Strategy for AI:** **NITI Aayog**'s 2018 National Strategy for AI introduced the vision of "AI for All" and promoted inclusive and socially beneficial artificial intelligence. The strategy emphasises safety, inclusivity, privacy, transparency, fairness and accountability as important foundations for responsible AI development.
- **IndiaAI Mission:** Launched in 2024, the IndiaAI Mission supports computing infrastructure, datasets, startups, skilling and innovation while including a "Safe and Trusted AI" pillar.
- **MeitY governance framework:** Ministry of Electronics and Information Technology guidelines promote a principle based techno-legal approach and use Seven Sutras for trustworthy, human centric, safe and responsible AI governance.

AI Ethics in World

International initiatives increasingly focus on human dignity, transparency, accountability, safety, human rights and responsible development of advanced AI systems.

- **UNESCO Recommendation:** Adopted in 2021, the UNESCO Recommendation on the Ethics of AI provides a global framework based on human dignity, human rights, transparency, accountability and sustainability.
- **EU Artificial Intelligence Act:** The 2024 law introduced a risk based regulatory model, prohibiting unacceptable risk applications such as social scoring and imposing stronger obligations on high risk AI systems.
- **Bletchley Park Declaration:** Signed by 28 countries in 2023, the Declaration recognised risks from frontier AI and promoted international cooperation on AI safety and responsible governance.
- **Global Partnership on AI:** GPAI is a multi stakeholder initiative supporting inclusive, responsible and human centric AI. India is a founding member and hosted the GPAI Summit in 2023.
- **Human centred governance:** Global frameworks increasingly support Human in the Loop systems, explainable AI, independent ethical assessments and safeguards for high risk applications involving life, liberty and dignity.

Rome Declaration Challenges

Turning the Rome Declaration for an Unarmed and Disarming Peace into effective international rules faces legal, technological, political and verification challenges across nuclear and artificial intelligence governance.

- **Lack of global consensus:** Major military powers may remain reluctant to accept binding restrictions on autonomous military capabilities, particularly when strategic competition is intensifying.
- **Verification difficulty:** AI software can be duplicated, modified, concealed and deployed remotely, making international verification substantially harder than monitoring physical nuclear materials.
- **Dual use technology:** The same AI technologies can support civilian research and military applications, making clear boundaries between beneficial innovation and dangerous deployment difficult to establish.
- **Rapid technological change:** Private sector AI development is progressing faster than multilateral treaty negotiations, creating a regulatory gap between technological capability and international governance.

Source: <https://vajiramandravi.com/current-affairs/rome-declaration/>

© Vajiram & Ravi. For personal use only.