



Text of Prime Minister's statement during the G7 Summit Session on Ensuring a Safe, Rapid and Efficient Rollout of AI

प्रविष्टि तिथि: 17 JUN 2026 9:20PM by PIB Delhi

Excellencies,

मैं इस महत्वपूर्ण विषय को हमारी चर्चा का भाग बनाने के लिए मेरे मित्र राष्ट्रपति मैक्रों का अभिनंदन करता हूँ। इसमें कोई शक नहीं है कि Artificial Intelligence मनुष्य द्वारा बनाई गई सबसे परिवर्तनकारी technologies में से एक है।

आज मानव जीवन का शायद ही कोई पहलू होगा, जिसे AI ने स्पर्श न किया हो। AI scientific रिसर्च को अभूतपूर्व गति दे रही है। Governance को अधिक effective और responsive बना रही है। स्वास्थ्य, शिक्षा, manufacturing जैसे क्षेत्रों को नई ताकत प्रदान कर रही है।

किन्तु, AI की वास्तविक कसौटी यह नहीं है कि हमारी मशीनें कितनी शक्तिशाली बनेंगी। इसकी असली कसौटी यह है कि सामान्य मानवी कितना empowered होगा। इस वर्ष भारत द्वारा आयोजित AI Impact Summit में हमने इसी सोच के साथ human-centric AI बनाने पर बल दिया। इस समिट में भारत ने अपना MANAV विज्ञान प्रस्तुत किया। यह vision AI में भारत के सभी प्रयासों को प्रेरित करता है।

हाल ही में “हिज़ होलीनेस द पोप” ने AI के विषय पर अपने पत्र में human values, inclusivity और meaningful human control को AI के विकास का आधार बनाने पर बल दिया है। भारत का MANAV vision और हिज़ होलीनेस का संदेश, दोनों एक ही मूल विचार को अभिव्यक्त करते हैं: टेक्नोलॉजी कितनी भी advanced क्यों न हो, उसके केंद्र में मानव ही रहना चाहिए।

Friends,

AI rollout में बच्चों के लिए safety सुनिश्चित करना अत्यंत आवश्यक है। AI बच्चों को उनकी अपनी भाषा में शिक्षा दे सकती है, उनकी creativity को बढ़ा सकती है, और learning को personalised बना सकती है। लेकिन safeguards के बिना यही टेक्नॉलजी उन्हें misinformation, deepfakes और exploitation के खतरे में डाल सकती है।

इन दोनों scenarios में फ़र्क टेक्नॉलजी का नहीं है। फ़र्क values का है, design का है, और governance का है। Digital space को हमें बच्चों के लिए learning का playground बनाना होगा, manipulation का tool नहीं।

Friends,

Frontier AI Models से Cyber Security के क्षेत्र में अभूतपूर्व संभावनाएं बन रही हैं। लेकिन Cyber Space में कोई भी देश तब तक पूरी तरह सुरक्षित नहीं हो सकता, जब तक सभी देश सुरक्षित न हों। इसलिए भारत ने हमेशा से Cyberspace को एक Global Public Good के रूप में देखा है। इसलिए इन महत्वपूर्ण AI Technologies तक पहुंच भी व्यापक और समावेशी होनी चाहिए। सभी लोकतांत्रिक देशों को ऐसे AI Models का access मिलना चाहिए, ताकि वे अपनी Critical Information Infrastructure की सुरक्षा कर सकें और बढ़ते Cyber Threats का सामना कर सकें।

Friends,

Safety, speed और efficiency की integrated approach पर आगे बढ़ने के लिए मैं कुछ सुझाव रखना चाहूंगा:

पहला, हमें safe-by-design AI systems को बढ़ावा देना चाहिए। Safety को बाद में जोड़ा गया feature नहीं, बल्कि design का मूल तत्व बनाना होगा।

दूसरा, AI deployment के लिए हमें common standards, testing frameworks और regulatory sandboxes विकसित करने चाहिए, ताकि innovation और governance साथ-साथ आगे बढ़ें। हमारे सामने सिविल एविएशन और मेरीटाइम ट्रांसपोर्ट ऐसे उदाहरण हैं जहाँ हमने global rules सफलतापूर्वक विकसित किये, और पूरे विश्व को इसका लाभ मिला।

तीसरा, deepfakes, misinformation और cyber fraud के विरुद्ध वैश्विक सहयोग को मजबूत करना होगा। हमें वॉटरमार्क्स जैसी टेक्नोलॉजीज़ को बढ़ावा देना चाहिए ताकि deepfakes से बचा जा सके।

चौथा, हमारा प्रयास होना चाहिए कि AI का लाभ ग्लोबल साउथ के सभी देशों तक पहुंचे, ताकि वह विभाजनकारी नहीं समावेशी शक्ति बने।

Friends,

AI के विषय में हमारी सोच और नीति स्पष्ट होनी चाहिए। AI must expand human potential, empower human choice, and protect human dignity. हम इस अत्यंत महत्वपूर्ण विषय पर सभी पार्टनर्स के साथ संवाद और सहयोग जारी रखेंगे।

बहुत-बहुत धन्यवाद।

MJPS/SS/ST

(रिलीज़ आईडी: 2274311) आगंतुक पटल : 512

इस विज्ञप्ति को इन भाषाओं में पढ़ें: हिन्दी , Assamese , Gujarati