

भारत कृत्रिम बुद्धिमत्ता (AI) की वैश्विक बहस को नई दिशा दे सकता है

UPSC मुख्य परीक्षा प्रासंगिकता

- GS पेपर II: वैश्विक शासन में भारत की भूमिका
- GS पेपर III: विज्ञान एवं प्रौद्योगिकी – विकास, अनुप्रयोग, डेटा संबंधी मुद्दे, AI और डिजिटल नैतिकता



समाचार में क्यों?

- ऐसे समय में जब दुनिया भर में कृत्रिम बुद्धिमत्ता (AI) को नियंत्रित करने के प्रयास तेज़ हो रहे हैं, भारत फरवरी 2026 में AI इम्पैक्ट समिट की मेज़बानी करेगा।
- भूराजनैतिक तनावों के बीच और AI गवर्नेंस पर वैश्विक सहमति के कमजोर पड़ने के माहौल में भारत के पास एक सुनहरा अवसर है:
- तकनीकी कूटनीति (Tech Diplomacy) में अपनी नेतृत्व भूमिका पुनः स्थापित करने का
- वैश्विक दक्षिण (Global South) की आकांक्षाओं का प्रतिनिधित्व करने का
- एक लोकतांत्रिक, समावेशी और नैतिक AI ढांचे का नेतृत्व करने का

प्रमुख वैश्विक AI सम्मेलन (Global AI Summits)

- **1. ब्लेवली पार्क, यूनाइटेड किंगडम (2023):**
- पहला AI सुरक्षा सम्मेलन आयोजित किया गया, जिसमें एडवांस्ड AI मॉडल्स के गवर्नेंस पर चर्चा हुई।
- **2. सियोल, दक्षिण कोरिया (2024):**
- ज़िम्मेदार एआई नवाचार और अंतर्राष्ट्रीय सहयोग पर ज़ोर दिया गया।
- **3. पेरिस, फ्रांस (2025):**
- वैश्विक एआई विनियमन को मानवाधिकारों, पारदर्शिता और नैतिक मानकों के अनुरूप बनाने के निरंतर प्रयास पर जोर दिया गया।

शिखर सम्मेलनों का महत्व: समावेशी और अंतरराष्ट्रीय दृष्टिकोण की आवश्यकता

- इन वैश्विक सम्मेलनों से यह स्पष्ट होता है कि कृत्रिम बुद्धिमत्ता के विकास और नियमन को समावेशी और अंतरराष्ट्रीय बनाया जाना आवश्यक है —
- जहाँ एक ओर नवाचार को बढ़ावा मिले, वहीं दूसरी ओर सुरक्षा और सार्वभौमिकता का संतुलन भी बना रहे।

लेकिन विभाजन अब भी बना हुआ है

- हालाँकि वैश्विक सहभागिता बढ़ रही है, फिर भी AI गवर्नेंस पर सर्वसम्मति अब तक नहीं बन पाई है।
- पेरिस AI सम्मेलन, 2025 में यह स्पष्ट हो गया कि प्रमुख शक्तियों के बीच मतभेद गहरे हो रहे हैं।
- सम्मेलन का उद्देश्य साझा मानदंड और सुरक्षा प्रोटोकॉल स्थापित करना था, लेकिन अंततः कोई एकीकृत घोषणा (Unified Declaration) नहीं हो सकी।
- अत्यधिक नियामक हस्तक्षेप और संप्रभुता के उल्लंघन की आशंका के कारण अमेरिका और यूनाइटेड किंगडम ने अंतिम घोषणापत्र का विरोध किया, वहीं चीन ने घोषणापत्र को समर्थन दिया — यह AI नियमन को लेकर दृष्टिकोणों की भिन्नता को दर्शाता है।
- यह सहमति की कमी वैश्विक AI शासन में गंभीर विखंडन की ओर इशारा करती है, जहाँ रणनीतिक, राजनीतिक और वैचारिक अंतर समन्वित प्रयासों को बाधित कर रहे हैं।

भारत की रणनीतिक बढ़त

- भारत वर्तमान वैश्विक AI परिदृश्य में एक निर्णायक मोड़ पर खड़ा है —
- यह केवल तकनीकी उपभोक्ता नहीं, बल्कि विकसित देशों और वैश्विक दक्षिण के बीच एक सेतु के रूप में स्थापित हो रहा है।
- भारत की विशिष्ट विशेषताएँ उसे वैश्विक AI मानदंडों को आकार देने में एक विश्वसनीय, लोकतांत्रिक और समावेशी आवाज़ बनाती हैं।

क्या बनाता है भारत को विशिष्ट?

- **1. लोकतांत्रिक वैधता :**
- विश्व के सबसे बड़े लोकतंत्र के रूप में भारत यह दिखाता है कि तकनीकी प्रगति को संवैधानिक अधिकारों और नागरिक स्वतंत्रताओं के अनुरूप कैसे रखा जा सकता है।
- **2. प्रौद्योगिकीय नेतृत्व :**
- आधार (दुनिया की सबसे बड़ी बायोमेट्रिक प्रणाली) और UPI (वैश्विक अग्रणी डिजिटल भुगतान प्रणाली) जैसे उदाहरण दिखाते हैं कि भारत जन-हित में AI-संलग्न अवसंरचना को बड़े पैमाने पर लागू करने में सक्षम है।
- **3. वैश्विक दक्षिण की आवाज़ (Voice of the Global South):**
- भारत बहुपक्षीय मंचों पर विकासशील देशों की आकांक्षाओं का प्रतिनिधित्व करता है — AI की समान पहुंच और न्यायसंगत शासन की वकालत करता है।
- **4. डिजिटल समावेशन (Digital Inclusion):**
- डिजिटल इंडिया जैसे कार्यक्रम यह दिखाते हैं कि भारत प्रौद्योगिकी को ग्रामीण और वंचित वर्गों तक पहुँचाने के लिए प्रतिबद्ध है।



जन भागीदारी और नीतिगत नवाचार

- भारत सरकार ने MyGov प्लेटफॉर्म के माध्यम से राष्ट्रीय स्तर पर परामर्श प्रक्रिया शुरू की है, जहाँ सुझाव आमंत्रित किए गए:
- छात्रों से
- स्टार्टअप्स से
- शिक्षाविदों और शोधकर्ताओं से
- नागरिक समाज संगठनों से
- यह भागीदारी आधारित दृष्टिकोण दर्शाता है कि भारत जन-केंद्रित और प्रासंगिक AI नीति ढांचा विकसित करने हेतु प्रतिबद्ध है।

वैश्विक AI बहस को दिशा देने के लिए भारत के 5 साहसिक विचार

- भारत के पास उद्देश्यपूर्ण, समावेशिता और व्यावहारिकता के साथ नेतृत्व करने का एक अनूठा अवसर है। अपने लोकतांत्रिक मूल्यों, तकनीकी गहराई और वैश्विक दक्षिण नेतृत्व से प्रेरित होकर, भारत ज़िम्मेदार और न्यायसंगत एआई विकास के लिए एक नया ढाँचा प्रस्तुत कर सकता

उद्देश्यपूर्ण संकल्प : महत्वपूर्ण बातों का आकलन

- सिर्फ सामान्य और अस्पष्ट घोषणाओं के बजाय, भारत को सरकारों, कंपनियों और अकादमिक संस्थानों से सटीक, मापनीय AI प्रतिबद्धताओं की माँग करनी चाहिए।

उदाहरण:

- एक टेक कंपनी द्वारा AI से होने वाली ऊर्जा खपत को कम करने की प्रतिज्ञा
- एक विश्वविद्यालय द्वारा ग्रामीण बालिकाओं को AI में प्रशिक्षित करने की योजना
- एक सरकार द्वारा AI-आधारित स्वास्थ्य उपकरणों को स्थानीय भाषाओं में अनुवादित कराने की व्यवस्था

पारदर्शिता तंत्र:

- इन सभी प्रतिज्ञाओं को एक सार्वजनिक वैश्विक डैशबोर्ड पर प्रकाशित किया जा सकता है, जहाँ प्रदर्शन रिपोर्ट कार्ड हर वर्ष अपडेट किया जाए — ताकि जवाबदेही सुनिश्चित हो।

वैश्विक दक्षिण को केंद्र में लाना

- भारत को वैश्विक एआई विमर्श में, विशेष रूप से कम प्रतिनिधित्व वाले क्षेत्रों में, समानता और समावेशन को बढ़ावा देना चाहिए। प्रमुख प्रस्तावों में शामिल हैं:-
- “AI for Billions Fund” की शुरुआत — जिसे विकास बैंकों और खाड़ी क्षेत्र के निवेशकों से समर्थन मिले
- Multilingual AI Challenge — 50 कम सेवा-प्राप्त भाषाओं के लिए AI उपकरण विकसित करने हेतु
- डेटा सेट, क्लाउड क्रेडिट, और वैश्विक AI फेलोशिप को ओपन एक्सेस के रूप में प्रचारित करना
- मूल संदेश: “प्रतिभा हर जगह है — सिर्फ कैलिफ़ोर्निया या बीजिंग में नहीं।”

साझा सुरक्षा मानक बनाना

- आज किसी देश में AI सुरक्षा के लिए एकीकृत चेकलिस्ट मौजूद नहीं है। भारत Global AI Safety Collaborative की स्थापना का प्रस्ताव दे सकता है:

प्रमुख उद्देश्य:

- Red Teaming प्रोटोकॉल और AI Stress Test के परिणाम साझा करना
- पूर्वाग्रह और मजबूती के मूल्यांकन प्रकाशित करना
- High-risk या High-impact AI के लिए मानक सीमाएँ तय करना
- ओपन-सोर्स सुरक्षा उपकरण इस मंच से विकसित किए जा सकते हैं — जो वैश्विक विश्वास और पारदर्शिता को बढ़ाएंगे।



संतुलित AI शासन: मध्य मार्ग की वकालत

- वैश्विक AI शासन की बहस वर्तमान में ध्रुवीकृत है। भारत मध्य मार्ग की व्यावहारिक और संयमित आवाज़ बन सकता है:
- क्षेत्र AI शासन का दृष्टिकोण
- अमेरिका अत्यधिक नियंत्रण का विरोध करता है
- यूरोपीय संघ कठोर AI अधिनियम लागू करता है
- चीन राज्य-नियंत्रण आधारित मॉडल अपनाता है
- भारत एक स्वैच्छिक “Frontier AI Code of Conduct” प्रस्तावित कर सकता है, जिसमें शामिल हों:
- 90 दिनों में Public Red-Team Disclosures
- एक तय सीमा से अधिक Compute Power Usage का खुलासा
- AI “Accident Hotline” की स्थापना — किसी गड़बड़ी या दुरुपयोग की स्थिति में वैश्विक समन्वित प्रतिक्रिया हेतु

विखंडन को रोकना: एक मंच, अनेक आवाज़ें

- AI गवर्नेंस के विखंडन से बचने के लिए भारत को समावेशी और सहयोगात्मक संवाद को प्रोत्साहन देना चाहिए — ताकि विकसित देशों और विकासशील देशों के बीच सेतु स्थापित हो।

रणनीतिक उद्देश्य:

- अमेरिका, यूरोपीय संघ, चीन, और वैश्विक दक्षिण के बीच सेतु निर्माण
- AI सम्मेलन का एजेंडा व्यापक और जन-केंद्रित बना रहे
- केवल प्रतिस्पर्धा नहीं, बल्कि साझा सार्वजनिक हितों पर फोकस किया जाए

निष्कर्ष:

- भारत का यह पांच सूत्रीय दृष्टिकोण न केवल वैश्विक AI नीति को संतुलित और न्यायसंगत बना सकता है, बल्कि AI को मानवता की भलाई के लिए दिशा देने वाला एक लोकतांत्रिक नेतृत्वकर्ता भी बना सकता है।
- दृष्टिकोण (Vision):
- AI पर एक एकीकृत वैश्विक दृष्टिकोण, जो विविध आवाज़ों को प्रतिबिंबित करे और तकनीकी रूप से खंडित भविष्य से बचाए।

भारत की आगे की राह (India's Path Ahead)

- भारत को एक केंद्रीकृत वैश्विक AI प्राधिकरण बनाने की ओर नहीं बढ़ना चाहिए।
- बल्कि उसकी शक्ति एक व्यावहारिक, समावेशी और सहयोगात्मक नेतृत्व प्रदान करने में है — जो उसके लोकतांत्रिक मूल्यों और विकास प्राथमिकताओं को प्रतिबिंबित करता है।

एक रणनीतिक और समावेशी दृष्टिकोण (A Strategic and Inclusive Approach)

- **1. मौजूदा पहलों का समन्वय (Integrate Existing Initiatives):**
- नई संस्थाएँ गढ़ने के बजाय, भारत वैश्विक और क्षेत्रीय प्रयासों को जोड़ने पर ध्यान दे सकता है — जिससे एक संगठित और सामंजस्यपूर्ण AI पारिस्थितिकी तंत्र बन सके।
- **2. वैश्विक बहुसंख्यक के साथ संसाधन साझा करना (Share AI Resources):**
- भारत अपने डिजिटल सार्वजनिक ढाँचे — जैसे ओपन-सोर्स टूल्स, डेटा सेट्स, और क्षमतावर्धन कार्यक्रम — को विकासशील देशों तक विस्तारित कर सकता है।
- **3. घरेलू मॉडलों को वैश्विक मंच पर प्रस्तुत करना (Showcase Homegrown Models):**
- आधार (डिजिटल पहचान) और UPI (डिजिटल भुगतान प्रणाली) जैसे मंच यह दिखाते हैं कि तकनीक को कैसे विस्तार, समानता और समावेशन के लिए प्रयोग में लाया जा सकता है — ये वैश्विक दक्षिण के लिए प्रतिकृति योग्य मॉडल हैं।
- **4. भाषा और डेटा विविधता को बढ़ावा देना (Promote Language and Data Diversity):**
- भारत बहुभाषी AI मॉडल्स, ओपन-एक्सेस डेटा रिपॉजिटरी, और AI के सुरक्षित व नैतिक विकास के लिए वैश्विक मानक स्थापित करने की वकालत कर सकता है — ताकि कोई क्षेत्र या समूह पीछे न छूटे।

भारत का व्यापक दृष्टिकोण (India's Larger Vision)

- भारत का उद्देश्य केवल सम्मेलन की मेज़बानी करना नहीं है, बल्कि डिजिटल युग में अपनी पहचान को पुनर्परिभाषित करना है — एक ऐसे राष्ट्र के रूप में:
- जो उत्तरदायी AI,
- प्रौद्योगिकीय समावेशन,
- और वैश्विक सहयोग का प्रवक्ता है।
- भारत न तो केवल एक रेग्युलेटर होगा, न ही एक मूक दर्शक, बल्कि एक सेतु-निर्माता (bridge-builder) बनेगा —
- जो तकनीकी नवाचार को सामाजिक न्याय, और वैश्विक शासन को स्थानीय सशक्तिकरण से जोड़ता है।

निष्कर्ष: भागीदारी को प्रगति में बदलना (Conclusion: Turning Participation into Progress)

- कृत्रिम बुद्धिमत्ता पर चल रही वैश्विक बहस केवल तकनीक की नहीं है — यह मूल्यों, न्याय और मानवता के भविष्य की भी बहस है।
- भारत के पास यह अवसर और ज़िम्मेदारी दोनों हैं कि वह इस परिवर्तन का नेतृत्व करे।

यदि भारत अपने:

- लोकतांत्रिक विश्वासनीयता,
- तकनीकी नवाचार, और
- समावेशी शासन मॉडल का रणनीतिक रूप से उपयोग करे, तो वह:
- जन-केंद्रित और विकासोन्मुखी AI ढाँचे का नेतृत्व कर सकता है
- वैश्विक उत्तर और दक्षिण के बीच संतुलनकारी सेतु बन सकता है
- वैश्विक AI शासन में नैतिक और नैतिकता आधारित आवाज़ बन सकता है
- वर्ष 2026 में प्रस्तावित AI समिट, यदि भारत द्वारा आयोजित होती है, तो यह केवल वैश्विक AI मानकों के निर्धारण का मोड़ नहीं होगी — बल्कि भारत की डिजिटल नेतृत्वकारी भूमिका को पुनः परिभाषित करने का क्षण भी बनेगी।

Result Mitra Mains Practice Questions (15 अंकों के लिए)

1. **Q1.** भारत के पास समावेशी और नैतिक कृत्रिम बुद्धिमत्ता पर वैश्विक बहस का नेतृत्व करने की साख और क्षमता दोनों हैं। भारत किन प्रमुख सिद्धांतों और रणनीतियों को अपनाकर वैश्विक AI शासन को दिशा दे सकता है? (15 अंक)
2. **Q2.** कृत्रिम बुद्धिमत्ता का भविष्य भूराजनीतिक प्रतिस्पर्धा नहीं, बल्कि वैश्विक सहयोग से निर्धारित होना चाहिए।

(वैकल्पिक विषय)
OPTIONAL SUBJECT
GEOGRAPHY
OPTIONAL

Fee - मात्र 6499 ₹

केवल 21 से
26 जून

Dr. Faiyaz Sir

OPTIONAL SUBJECT
वैकल्पिक विषय
PSIR

Fee - मात्र 6999 ₹

केवल 01 से
06 जुलाई

Dr. Faiyaz Sir