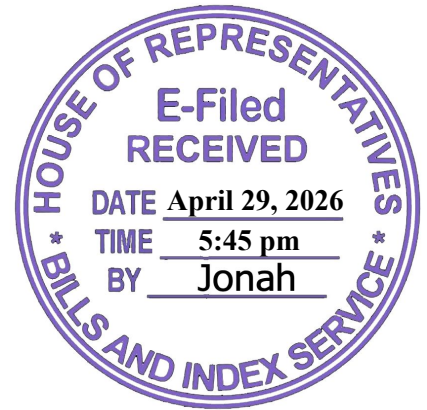




Republic of the Philippines
HOUSE OF REPRESENTATIVES
Quezon City

TWENTIETH CONGRESS
First Regular Session



HOUSE RESOLUTION NO. 971

Introduced by CIBAC Party-List Representative
EDUARDO “BRO. EDDIE” C. VILLANUEVA

RESOLUTION

DIRECTING THE COMMITTEE ON SCIENCE AND TECHNOLOGY AND OTHER APPROPRIATE COMMITTEE/S OF THE HOUSE OF REPRESENTATIVES TO CONDUCT AN INQUIRY, IN AID OF LEGISLATION, ON THE NATIONAL PREPAREDNESS FOR RISKS ASSOCIATED WITH ARTIFICIAL INTELLIGENCE (AI), INCLUDING SYSTEMIC, ECONOMIC, HEALTH, SECURITY, AND CATASTROPHIC RISKS, WITH THE END IN VIEW OF DEVELOPING A COMPREHENSIVE NATIONAL FRAMEWORK FOR AI RISK GOVERNANCE

WHEREAS, Artificial Intelligence (AI) refers to machine-based systems designed to perform tasks that normally require human intelligence, including learning from data, recognizing patterns, making predictions or decisions, and generating content, which may affect real or virtual environments^{1, 2};

WHEREAS, AI systems are rapidly transforming economic activity, governance, national security, and social life, and are increasingly integrated into critical sectors such as education, labor markets, agriculture, public administration, finance, and digital infrastructure;

WHEREAS, modern machine learning systems derive their capabilities from the interaction of data, algorithms, and computing power, commonly referred to as the AI triad, which constitutes the foundational architecture of contemporary AI development and performance³;

WHEREAS, the global distribution of these resources remains highly uneven, with a limited number of technologically advanced states and private corporations controlling the

¹ 2019 May. *Recommendation of the Council on Artificial Intelligence*. Organisation for Economic Co-operation and Development. Available at <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>. Accessed on 16 March 2026.

² 2023 January. *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. National Institute of Standards and Technology, U.S. Department of Commerce. Available at <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>. Accessed on 16 March 2026.

³ 2020 August. *The AI Triad and What It Means for National Security*. Center for Security and Emerging Technology. Available at <https://cset.georgetown.edu/wp-content/uploads/CSET-AI-Triad-Report.pdf>. Accessed on 11 March 2026

majority of high-performance computing infrastructure, training datasets, and frontier AI research capacity;

WHEREAS, the Philippines currently functions primarily as a consumer of AI technologies developed abroad, resulting in structural asymmetries in the nation's ability to evaluate, regulate, or constrain the risks posed by such technologies;

WHEREAS, an assessment of national preparedness across the following categories of AI risks is imperative to be done:

1. Systemic Societal Risks - Include misinformation and disinformation generation, manipulation of public discourse, erosion of trust in digital information ecosystems, and algorithmic bias,.
2. Economic and Labor Disruption Risks - Include workforce displacement, labor market restructuring, and economic concentration associated with automation and AI deployment.
3. Technological and Strategic Dependency Risks - Include reliance on foreign AI technologies, limited national access to computing infrastructure, and the implications of technological dependency for national sovereignty.
4. Security and Infrastructure Risks - Include AI-enabled cyberattacks, vulnerabilities in automated systems, and national security implications of advanced AI capabilities.
5. Health Risks – Include clinical decision-tools, diagnostic AI systems, public health analytics, and digital health platforms.
6. Catastrophic and Emerging Risks - Include the misuse of AI systems in biological, cyber, or information warfare contexts, and the potential emergence of advanced autonomous AI agents.

WHEREAS, these risks include the generation of misinformation, manipulation of democratic processes, algorithmic discrimination, cybersecurity vulnerabilities, and malicious misuse of AI systems;

WHEREAS, research on AI safety reveals that safety is not a property of AI models, thus, it must instead be evaluated with caution;⁴

WHEREAS, the government should have the capacity to conduct AI safety evaluations in order to effectively regulate and mitigate AI risks⁵;

WHEREAS, global regulatory initiatives such as the European Union Artificial Intelligence Act have begun to establish risk-based governance frameworks that classify AI systems according to their societal risks and impose stronger regulatory obligations on high-risk AI applications used in sectors such as education, employment, critical infrastructure, and law enforcement⁶;

WHEREAS, the integration of AI technologies across key sectors of Philippine society presents both opportunities for innovation and risks related to labor disruption, technological dependency, national security vulnerabilities, and systemic governance challenges;

⁴ 2025, March 13. *AI Safety is not a model property*. AI as Normal Tech. Available at <https://www.normaltech.ai/p/ai-safety-is-not-a-model-property>. Accessed on 11 March 2026.

⁵ 2025, May 28. *AI Safety Evaluations: An Explainer*. Center for Security and Emerging Technology. Available at <https://cset.georgetown.edu/article/ai-safety-evaluations-an-explainer>. Accessed on 11 March 2026.

⁶ 2024, February 27. *High-level summary of the [EU] AI Act*. EU Artificial Intelligence Act. Available at <https://artificialintelligenceact.eu/high-level-summary/>. Accessed on 11 March 2026.

AI risks on labor sector

WHEREAS, artificial intelligence technologies are already affecting labor markets worldwide, with companies reporting job reductions linked to automation and AI adoption, and studies identifying occupations such as customer service, programming, and data entry as among the roles most exposed to AI-driven automation^{7, 8};

WHEREAS, studies indicate that millions of workers globally are employed in occupations highly exposed to artificial intelligence technologies, highlighting the potential for AI systems to automate or significantly reshape routine and cognitive work tasks across multiple industries⁹;

AI risks on education

WHEREAS, artificial intelligence technologies are increasingly being integrated into education systems through tools such as automated tutoring systems, personalized learning platforms, and AI-assisted feedback mechanisms that significantly influence teaching and learning processes¹⁰;

WHEREAS, researchers and educators have also raised concerns that the growing use of AI tools in classrooms may introduce new challenges including academic dishonesty, reduced student engagement, and broader questions regarding the role of AI in shaping learning outcomes¹¹;

Catastrophic risks of AI

WHEREAS, emerging research further warns that advanced AI systems may introduce catastrophic risks, including the potential misuse of AI systems in biological research, large-scale cyber operations, and automated persuasion systems capable of influencing democratic institutions¹²;

WHEREAS, scholars and researchers emphasize the importance of AI alignment, referring to the challenge of ensuring that AI systems reliably pursue goals consistent with human intentions, implying that misaligned AI systems have the capacity to perform unintended objectives with potentially harmful consequences¹³;

⁷ 2026, February 27. *More companies are pointing to AI as the lay off employees*. CBS News. Available at <https://www.cbsnews.com/news/ai-layoffs-2026-artificial-intelligence-amazon-pinterest/>. Accessed on 11 March 2026.

⁸ 2026, March 06. *Programming, customer service, data entry roles among jobs most at risk from AI: Anthropic*. The Economic Times. Available at <https://economictimes.indiatimes.com/tech/artificial-intelligence/programming-customer-service-data-entry-roles-among-jobs-most-at-risk-from-ai-anthropic/articleshow/129158865.cms?from=mdr>. Accessed on 11 March. 2026.

⁹ 2026 February 19. *6.1 Million Workers Have the Most to Lose From AI*. Investopedia. Available at <https://www.investopedia.com/6-1-million-workers-have-the-most-to-lose-from-ai-11909256>. Accessed on 11 March 2026.

¹⁰ 2024, February 5. *Impact of Artificial Intelligence in Education*. eSchool News. Available at <https://www.eschoolnews.com/digital-learning/2024/02/05/impact-of-artificial-intelligence-in-education/>. Accessed on 11 March 2026.

¹¹ 2025, October 08. *Rising Use of AI in Schools Comes with Big Downsides for Students*. Education Week. Available at <https://www.edweek.org/technology/rising-use-of-ai-in-schools-comes-with-big-downsides-for-students/2025/10>. Accessed on 11 March 2026.

¹² 2025, August 26. *From Principles to Rules: A Regulatory Approach for Frontier AI*. Cornell University. Available at <https://arxiv.org/abs/2407.07300>. Accessed on 11 March 2026.

¹³ 2024, March 02. *What is AI Alignment*. Bluedot Impact. Available at <https://blog.bluedot.org/p/what-is-ai-alignment>. Accessed on 11 March 2026.

WHEREAS, the absence of a coordinated national framework for AI risk governance may leave critical sectors of Philippine society exposed to technological risks;

WHEREAS, a comprehensive national assessment of AI risks across sectors is necessary to inform legislation that will ensure the safe, responsible, and sovereign deployment of AI technologies within the Philippines;

Examples of AI Risk Governance

WHEREAS, governments around the world have begun developing institutional frameworks for the governance and risk management of artificial intelligence, recognizing the need to assess and mitigate potential societal, economic, and security risks associated with increasingly capable AI systems;

WHEREAS, for example, the United States National Institute of Standards and Technology (NIST) has developed the Artificial Intelligence Risk Management Framework (AI RMF) to guide governments, organizations, and industries in identifying, assessing, and managing risks associated with the design, development, deployment, and use of AI systems¹⁴;

WHEREAS, the European Union has enacted the Artificial Intelligence Act, the world's first comprehensive legal framework regulating artificial intelligence, which adopts a risk-based approach that classifies AI systems according to levels of societal risk and imposes stricter regulatory obligations on high-risk AI applications used in sectors such as education, employment, critical infrastructure, and law enforcement;¹⁵

WHEREAS, the United Kingdom has established the AI Safety Institute, a government-backed institution dedicated to evaluating the safety and risks of frontier artificial intelligence systems and supporting the development of policy mechanisms for AI governance¹⁶;

WHEREAS, several countries and international partners are participating in the development of a global ecosystem of AI Safety Institutes, designed to evaluate advanced AI systems and coordinate international approaches to AI safety and governance¹⁷;

WHEREAS, these international developments demonstrate a growing global consensus that governments must proactively assess and manage the risks associated with artificial intelligence technologies through structured governance frameworks, institutional oversight mechanisms, and cross-sector policy coordination;

NOW, THEREFORE, BE IT RESOLVED BY THE HOUSE OF REPRESENTATIVES, That the Committee on Science and Technology, jointly with other appropriate committees of the House of Representatives, conduct an inquiry, in aid of legislation, into the

¹⁴ 2023, January. *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. National Institute of Standard and Technology, U.S. Department of Commerce. Available at <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>.

¹⁵ 2024. *EU AI Act*. European Union. Available at <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>. Accessed on 11 March 2026.

¹⁶ 2024, January 17. *AI Safety Institute: overview*. Department for Science, Innovation and Technology and AI Safety Institute, United Kingdom. Available at <https://www.gov.uk/government/publications/ai-safety-institute-overview>. Accessed on 11 March 2026.

¹⁷ 2024, July 29. *AI Safety Institutes: Can countries meet the challenge?* Organisation for Economic Co-operation and Development. Available at <https://oecd.ai/en/work/ai-safety-institutes-challenge>. Accessed on 11 March 2026.

preparedness of the Philippine government to assess, monitor, and manage the risks associated with artificial intelligence systems across key national sectors.

RESOLVED FURTHER, That all concerned government agencies, instrumentalities, and government-owned or controlled corporations shall be required to present their respective artificial intelligence-related initiatives, innovations, and digital systems, if any, currently deployed, under development, or in pilot stages within their respective sectors;

RESOLVED FURTHER, That such presentations may include, but not be limited to:

1. **Description of AI Innovations and Systems**
 - Specific AI tools, platforms, or systems currently in use or being developed
 - Objectives and intended functions of such systems
2. **Stage of Development and Deployment**
 - Whether such systems are in pilot testing, partial deployment, or full operational use
3. **Observed Benefits and Outcomes**
 - Efficiency gains, cost reductions, or service improvements attributed to the use of AI systems
4. **Identified Risks and Challenges**
 - Issues encountered, including bias, system errors, data privacy concerns, or operational limitations
5. **Safeguards and Oversight Measures**
 - Mechanisms in place to ensure accountability, transparency, and safe deployment of AI systems

RESOLVED FURTHER, That the inquiry shall examine the current state of AI deployment, governance, and risk preparedness across major sectors, including but not limited to, education, health, labor and employment, agriculture and food systems, science and technology development, public administration and governance, national security and defense, financial systems and economic regulation, and digital infrastructure and communications. All concerned government agencies, instrumentalities, research institutions, and relevant stakeholders shall be invited to participate in the inquiry.

RESOLVED FINALLY, That the findings of the inquiry shall guide the development of legislation establishing a national framework for AI risk governance, including, national AI safety standards, mandatory AI risk assessments for government deployment, sector-specific AI regulatory frameworks, institutional mechanisms for AI safety evaluation and oversight, and strategies to strengthen national technological capacity and reduce dependency on foreign AI systems.

Adopted,


EDUARDO "BRO. EDDIE" C. VILLANUEVA