

ICITAIDS

International Conference on Interdisciplinary Trends on Artificial Intelligence & Data Science (ICITAIDS - 2026)

organized by:



International Scientific Society
Sharing Knowledge & Wisdom

(An International community for Independent & academic Research Scholars)



ISSN : 2583-2646

In Association With:



ESP Journal of Engineering & Technology Advancements
International Journal, Scholarly Peer-Reviewed and Published Globally

International Conference on Interdisciplinary Trends on Artificial Intelligence & Data Science (ICITAIDS - 2026)

Proceedings of

ICITAIDS - 2026

07th January 2026

Organized by



International Scientific Society
Sharing Knowledge & Wisdom
(An International community for Independent & academic Research Scholars)

In Association With



ESP Journal of Engineering & Technology Advancements
International Journal, Scholarly Peer-Reviewed and Published Globally

VVS Arcade, 18/1, Puthur High Road,
Opposite to Aruna Theater,
Tiruchirappalli - 620017

About the Conference

The International Conference on Interdisciplinary Trends on Artificial Intelligence & Data Science (ICITAIDS–2026) is a premier international academic conference that aims to bring together researchers, academicians, scientists, industry professionals, and students from around the world to discuss the latest developments, innovations, and challenges in the fields of Artificial Intelligence (AI) and Data Science. The conference provides a multidisciplinary platform where experts from various domains can exchange knowledge, present high-quality research, and establish collaborations that contribute to technological advancement and societal development. The conference theme emphasizes the integration of AI and Data Science with interdisciplinary research, encouraging innovative solutions to real-world problems across diverse sectors.

ICITAIDS–2026 focuses on promoting research that combines Artificial Intelligence, Machine Learning, Deep Learning, Data Analytics, Computer Vision, Natural Language Processing, Robotics, Internet of Things (IoT), Cloud Computing, Big Data Analytics, Cybersecurity, Blockchain, Intelligent Systems, and other emerging technologies. The conference recognizes that modern technological challenges require interdisciplinary approaches, where AI and Data Science can be effectively integrated with healthcare, agriculture, education, finance, transportation, manufacturing, environmental science, and smart city applications. By encouraging collaboration across multiple disciplines, the conference seeks to accelerate innovation and foster sustainable technological growth. (icaids.com -)

One of the primary objectives of ICITAIDS–2026 is to provide a global platform for presenting original research findings and innovative ideas. Participants have the opportunity to present their research papers through technical sessions, keynote lectures, invited talks, panel discussions, workshops, and poster presentations. These sessions enable researchers to receive valuable feedback from experienced experts while promoting academic discussions on emerging technologies and future research directions. The conference also encourages interaction between academia and industry, enabling participants to understand practical challenges and identify opportunities for collaborative research and technology transfer.

The conference welcomes contributions from universities, research institutions, government organizations, startups, and leading technology companies. Distinguished keynote speakers, internationally recognized researchers, and industry experts share their experiences, research achievements, and insights into future technological developments. These expert sessions provide participants with valuable exposure to current research trends, practical applications, and emerging innovations in Artificial Intelligence and Data Science. The conference also serves as an excellent networking platform, allowing participants to establish professional relationships, exchange ideas, and initiate future research collaborations with experts from different countries and disciplines.

ICITAIDS–2026 covers a broad range of research topics, including Artificial Intelligence, Machine Learning, Deep Learning, Data Mining, Predictive Analytics, Computer Vision, Natural Language Processing, Edge Computing, Federated Learning, Explainable AI, Generative AI, Reinforcement Learning, Smart Healthcare, Intelligent Transportation Systems, Financial Technologies, Smart Agriculture, Industry 5.0, Sustainable Computing, Human-Computer Interaction, Ethical AI, AI Governance, Digital Transformation, and Intelligent Decision Support Systems. These topics reflect the growing importance of AI-driven technologies in solving complex societal and industrial challenges.

In addition to technical knowledge sharing, the conference promotes innovation, entrepreneurship, and interdisciplinary research by encouraging participants to explore new applications of Artificial Intelligence and Data Science in solving global challenges. Young researchers, postgraduate students, and doctoral scholars benefit from exposure to cutting-edge research methodologies, publication opportunities, and interactions with internationally recognized experts. Such engagement enhances research quality, supports professional development, and inspires future innovations in intelligent technologies.

Overall, the International Conference on Interdisciplinary Trends on Artificial Intelligence & Data Science (ICITAIDS–2026) serves as an important international forum dedicated to advancing research, encouraging interdisciplinary collaboration, and promoting technological innovation. By bringing together academic researchers, industry professionals, policymakers, and students under a common platform, the conference contributes significantly to the advancement of Artificial Intelligence and Data Science while addressing the evolving needs of modern society through collaborative research and innovative technological solutions.

Keynote Speakers / Invited Speakers

Dr. Abudl Hannan Abdul Mannan Shaikh
Asst. Professor, Department of Computer Science and Information Technology,
Al-Baha University, Kingdom of Saudi Arabia

Abhishek Baer
Customer Resident Quality Engineer & Independent Researcher, USA.

Aakash Gadh
FOD Control Corporation, USA.

Lakshya Khurana
Independent Researcher, USA.

Yashika Shankheshwaria
Sales IT Data Analyst, USA.

Prakash Velusamy
Principal Software Development Engineer, USA.

Venkateswara Varma Srivatsavaya
Principal Solutions Architect, Cloudera USA.

Upendra Kanuru
DevOps Engineer & Independent Researcher , USA.

Sai Nikhil Donthi
Specialist Software Engineering.

Hemant Soni
Digital Transformation & Business Optimization Leader.

Ameya Kokate
Senior Data & Analytics Engineer | AI/ML Researcher, USA.

Senthil Nathan
US Global Payment Platform Product Lead, J.P. Morgan, USA.

Pankaj Kumar
Ex Technical Lead TCS, USA.

Advisory Committee

Dr. Jayasudha Yedalla
Software Engineer, Independent Researcher, USA.

Dr.A.Karunamurthy
Associate Professor, Department of MCA Sri Manakula Vinayagar Engineering College (Autonomous)
Puducherry, India.

Mr. Ashutosh Agarwal
Senior Analyst in Data Engineering and Independent Researcher, USA.

Dr. J Sirisha Devi
Associate Professor, Department of Computer Science & Engineering, Institute of Aeronautical Engineering (IARE),Hyderabad, India.

Dr. Ts. Keith Chong Peng Lean
Faculty of Engineering and Technology, Multimedia University, Melaka, Malaysia.

Sriram Ghanta
Designation Senior Java Full Stack Developer,research scholar,USA.

Dr. Victor Sunday Aigbodon
Professor, Department of Metallurgical And Materials Engineering,
Faculty of Engineering, University of Nigeria, Nsukka Nigeria.

Judges / Peer Reviewers

Suraj Balaso Desai
Data Scientist, Digg Inc.

Prassanna Rao Rajgopal
Cybersecurity Leader & Independent Researcher, USA.

Raghava Chellu
Support Engineer @ Equifax Inc & Independent Researcher, USA.

Ranjith Kumar Ramakrishnan
Senior Software Developer & Independent Researcher, USA.

Srinivasa Rao Seetala
Senior Data Architect , USA.

Dr. Kumar Debasis
Associate Professor Grade 1 Department of Computer Science & Engg. VIT-AP University, Andhra Pradesh,
India.

CONTENTS

S.No	Title / Author Name	Page No
1.	World Foundation Models for Autonomous Physical Intelligence in Next-Generation Smart Systems Dr. Arjun Menon¹, Priya Sharma	1
2.	Self-Evolving Agentic AI Through Recursive Test-Time Learning for Autonomous Decision Intelligence Dr. Michael Anderson¹, Rahul Krishnan², Dr. Sneha Iyer³	2
3.	Memory-Augmented Large Reasoning Models for Long-Horizon Autonomous Task Planning Dr. Kavin Raj¹, Emily Johnson²	3
4.	Neuro-Symbolic World Models for Explainable Scientific Discovery and Autonomous Knowledge Reasoning Dr. Ethan Walker¹, Divya Nair²	4
5.	Spatial Intelligence Framework for Real-Time Human-AI Collaboration in Physical Environments Dr. Harish Kumar¹, John Peterson², Ananya Rao³	5
6.	Universal AI Operating Systems for Coordinating Autonomous Multi-Agent Ecosystems Dr. Sophia Carter¹, Vikram Singh²	6
7.	Hybrid Quantum-Agentic Computing Architecture for High-Performance Scientific Problem Solving Dr. Naveen Babu¹, Keerthana Raj²	7
8.	Self-Verifying Foundation Models for Reliable Autonomous Research and Decision Making Dr. William Harris¹, Riya Patel²	8
9.	Cognitive Digital Organisms: Self-Adaptive Artificial Intelligence for Autonomous Computing Networks Dr. Suresh Narayanan¹, Jennifer Brown², Dr. Daniel Lee³	9
10.	Embodied Foundation Models for Intelligent Human-Robot Collaboration in Industry 6.0 Dr. Rohit Verma¹, Olivia Martinez²	10
11.	Synthetic Universe Generation for Training Autonomous Physical Artificial Intelligence Systems Dr. Akash Reddy¹, Sophia Wilson²	11
12.	Hierarchical Multi-Agent Intelligence Framework for Autonomous Enterprise Workflow Optimization Dr. Benjamin Clark¹, Nithya Lakshmi²	12
13.	Adaptive Neural Memory Networks for Lifelong Continual Learning in Foundation Models Dr. Karthik Srinivasan¹, Grace Thomas², Aravind Raj³	13
14.	Trust-Calibrated Agentic Artificial Intelligence for Human-Centered Decision Support Systems Dr. Rachel Moore¹, Pranav Nair²	14
15.	Physics-Grounded World Models for Autonomous Navigation in Unknown Dynamic Environments Dr. Lakshmi Prasad¹, Jacob Miller²	15
16.	Energy-Adaptive Green Foundation Models for Sustainable Artificial Intelligence Infrastructure Dr. David Wilson¹, Meena Krishnan², Dr. Vivek Sharma³	16
17.	Test-Time Knowledge Editing Framework for Continually Learning Large Language Models Dr. Arun Prakash¹, Jessica Taylor²	17
18.	Collective Swarm Intelligence Using Cooperative Agentic AI for Distributed Autonomous Systems Dr. Christopher Evans¹, Aishwarya Menon²	18
19.	Cross-Reality Artificial Intelligence Integrating Physical, Digital, and Spatial Computing Ecosystems Dr. Senthil Kumar¹, Ethan Roberts²	19
20.	Autonomous AI Scientists for Scientific Hypothesis Generation and Experimental Validation Dr. Abigail Scott¹, Rahul Desai², Dr. Kavya Raman³	20
21.	Self-Healing Artificial Intelligence Infrastructure Using Autonomous Multi-Agent Orchestration Dr. Mohan Raj¹, Sarah Thompson²	21
22.	Foundation Models for Universal Scientific Simulation Across Interdisciplinary Research Domains Dr. Noah Mitchell¹, Pooja Kapoor²	22
23.	Causal World Models for Autonomous Decision Intelligence in Complex Dynamic Systems Dr. Gokul Krishna¹, Emily Davis², Arjun Patel³	23
24.	Xeromorphic Argentic Computing for Ultra-Efficient Next-Generation Artificial Intelligence Dr. Robert Green¹, Harini Suresh²	24

25.	World Model-Based Autonomous Cyber Defense Using Intelligent Multi-Agent Threat Reasoning <i>Dr. Vinay Kumar¹, Laura White²</i>	25
26.	Large Reasoning Models for Explainable Scientific Computing Beyond Conventional Large Language Models <i>Dr. Andrew Collins¹, Deepika Iyer²</i>	26
27.	AI-Native Autonomous Computing Platforms for Future Intelligent Digital Ecosystems <i>Dr. Praveen Natarajan¹, Kevin Brooks²</i>	27
28.	Human-Centered Artificial General Intelligence Framework for Safe and Responsible Autonomous Systems <i>Dr. Matthew Robinson¹, Shruthi Narayanan², Dr. Rajesh Kumar³</i>	28

World Foundation Models for Autonomous Physical Intelligence in Next-Generation Smart Systems

Dr. Arjun Menon¹, Priya Sharma

¹Associate Professor, Department of Computer Science and Engineering, Indian Institute of Technology Bombay, Mumbai, India

²AI Research Engineer, Google LLC, Mountain View, California, USA

Abstract: The rapid advancement of artificial intelligence has accelerated the development of autonomous systems capable of interacting with complex physical environments. However, conventional AI models often depend on task-specific training and lack the generalized reasoning required to adapt to dynamic real-world scenarios. This research proposes a World Foundation Model (WFM) framework that enables autonomous physical intelligence through comprehensive environmental understanding, predictive reasoning, and adaptive decision-making. The proposed architecture integrates multimodal perception, spatial-temporal reasoning, causal inference, and continual learning to construct a unified representation of the physical world. By combining vision, language, sensor streams, and contextual information, the model develops an internal world representation that supports long-term planning, autonomous navigation, object interaction, and real-time problem solving. A hierarchical memory mechanism continuously updates environmental knowledge while reinforcement learning and self-supervised learning improve adaptability in previously unseen situations. The framework is designed for deployment across next-generation smart systems, including intelligent manufacturing, autonomous transportation, healthcare robotics, precision agriculture, disaster response, and smart infrastructure management. Performance evaluation considers perception accuracy, reasoning capability, planning efficiency, adaptability, energy consumption, and safety compliance using both simulated and real-world environments. Experimental analysis demonstrates that the proposed approach significantly improves decision consistency, environmental awareness, operational reliability, and autonomous task completion compared with traditional deep learning architectures. The research also addresses challenges related to scalability, computational efficiency, trustworthiness, and ethical deployment of autonomous intelligence. The proposed World Foundation Model represents a significant step toward the realization of highly adaptive, explainable, and general-purpose physical intelligence capable of supporting the future generation of intelligent cyber-physical ecosystems.

Keywords: World Foundation Models, Physical Intelligence, Autonomous Systems, Smart Systems, Multimodal Learning, Spatial Reasoning, Continual Learning, Reinforcement Learning, Intelligent Automation, Cyber-Physical Systems

Self-Evolving Agentic AI Through Recursive Test-Time Learning for Autonomous Decision Intelligence

Dr. Michael Anderson¹, Rahul Krishnan², Dr. Sneha Iyer³

¹Professor, Department of Computer Science, Stanford University, California, USA

²Research Scholar, Department of Artificial Intelligence, Anna University, Chennai, India

³Associate Professor, Department of Information Technology, Vellore Institute of Technology, Vellore, India

Abstract: Artificial intelligence is rapidly evolving from task-oriented automation toward autonomous agents capable of continuous learning and intelligent decision-making. Despite significant advancements, most existing AI models remain static after deployment and require periodic retraining to adapt to changing environments. This research introduces a Self-Evolving Agentic AI framework based on Recursive Test-Time Learning (RTTL) that enables autonomous systems to refine their reasoning, update internal knowledge, and optimize decisions during real-world operation without complete model retraining. The proposed architecture combines large reasoning models, dynamic memory management, meta-learning, confidence-aware inference, and recursive feedback optimization to support continual adaptation across diverse environments. During inference, the agent evaluates its own predictions, identifies uncertainty, incorporates newly observed information, and recursively improves **future** decision-making while preserving previously acquired knowledge. A hierarchical cognitive controller coordinates planning, reasoning, execution, and self-correction to enhance long-term autonomous behavior. The framework is evaluated across intelligent transportation, industrial automation, healthcare decision support, cybersecurity threat response, financial forecasting, and smart infrastructure management. Performance metrics include decision accuracy, adaptation speed, reasoning consistency, computational efficiency, knowledge retention, uncertainty estimation, and autonomous task completion rate. Experimental evaluation demonstrates that recursive test-time learning significantly improves system robustness, reduces catastrophic forgetting, increases environmental adaptability, and minimizes human intervention compared with conventional fine-tuning and static inference approaches. The proposed framework also incorporates safety constraints, explainability mechanisms, and ethical governance to ensure reliable autonomous operation in mission-critical applications. This research establishes a scalable foundation for future self-improving intelligent agents capable of lifelong learning, adaptive reasoning, and autonomous decision intelligence in complex real-world ecosystems.

Keywords: Agentic AI, Recursive Test-Time Learning, Autonomous Decision Intelligence, Lifelong Learning, Adaptive Reasoning, Large Reasoning Models, Meta-Learning, Autonomous Agents, Explainable AI, Intelligent Decision Systems

Memory-Augmented Large Reasoning Models for Long-Horizon Autonomous Task Planning

Dr. Kavin Raj¹, Emily Johnson²

¹Associate Professor, Department of Artificial Intelligence, SRM Institute of Science and Technology, Chennai, India

²Senior Machine Learning Engineer, Microsoft Corporation, Redmond, Washington, USA

Abstract: The emergence of Large Reasoning Models (LRMs) has significantly improved artificial intelligence capabilities in complex problem-solving and logical inference. However, existing reasoning models face limitations in maintaining long-term contextual understanding, preserving historical knowledge, and executing extended sequences of interconnected tasks. This research proposes a Memory-Augmented Large Reasoning Model (MA-LRM) designed to support long-horizon autonomous task planning through hierarchical memory integration, contextual reasoning, and adaptive knowledge retrieval. The proposed framework combines semantic memory, episodic memory, working memory, and external knowledge repositories to enable intelligent agents to recall previous experiences, interpret evolving environmental contexts, and formulate optimized multi-stage execution strategies. A dynamic memory controller continuously prioritizes, updates, and compresses relevant information while minimizing redundancy and preventing catastrophic forgetting. The reasoning engine incorporates causal inference, symbolic reasoning, and reinforcement learning to improve planning accuracy and decision consistency across complex operational environments. The framework is evaluated using applications in autonomous robotics, smart manufacturing, healthcare workflow automation, intelligent logistics, financial planning, scientific research assistance, and disaster management. Performance evaluation includes planning success rate, reasoning accuracy, memory retrieval efficiency, computational overhead, adaptability, knowledge retention, execution latency, and long-term task completion performance. Experimental analysis demonstrates that integrating hierarchical memory significantly enhances contextual understanding, reduces planning errors, improves sequential decision-making, and increases operational reliability compared with conventional large language and reasoning models. The proposed architecture further incorporates explainable reasoning mechanisms, trust-aware memory validation, and secure knowledge management to support transparent and dependable autonomous operation. This research contributes a scalable cognitive architecture capable of sustaining intelligent long-term planning, continuous learning, and autonomous reasoning, thereby advancing the development of next-generation artificial intelligence systems capable of solving complex real-world problems with minimal human intervention.

Keywords: Large Reasoning Models, Memory-Augmented AI, Long-Horizon Planning, Autonomous Task Planning, Hierarchical Memory, Causal Reasoning, Reinforcement Learning, Continual Learning, Intelligent Agents, Cognitive Artificial Intelligence

Neuro-Symbolic World Models for Explainable Scientific Discovery and Autonomous Knowledge Reasoning

Dr. Ethan Walker¹, Divya Nair²

¹Professor, School of Engineering, Massachusetts Institute of Technology, Cambridge, Massachusetts, USA

²Research Scholar, Department of Computer Science, National Institute of Technology Calicut, India

Abstract: The increasing complexity of scientific research demands artificial intelligence systems capable of not only identifying hidden patterns in large-scale datasets but also providing transparent reasoning that supports trustworthy knowledge discovery. Although deep neural networks demonstrate remarkable predictive performance, they often operate as black-box models with limited interpretability, while symbolic reasoning systems struggle to process high-dimensional and unstructured information. This research proposes a Neuro-Symbolic World Model (NSWM) that integrates deep learning, symbolic reasoning, causal inference, and world modeling to facilitate explainable scientific discovery and autonomous knowledge reasoning. The proposed architecture combines multimodal perception, knowledge graphs, logical rule engines, semantic memory, and dynamic reasoning modules to construct an evolving representation of complex scientific environments. The framework continuously acquires knowledge from structured and unstructured sources, validates hypotheses through causal reasoning, and generates transparent decision pathways that can be verified by domain experts. A hierarchical reasoning engine enables autonomous hypothesis generation, evidence evaluation, contradiction detection, and iterative knowledge refinement while maintaining consistency across multiple scientific domains. The proposed model is evaluated using applications in biomedical research, climate science, materials engineering, pharmaceutical discovery, astrophysics, and computational chemistry. Performance evaluation considers reasoning accuracy, hypothesis validation rate, explainability score, knowledge consistency, computational efficiency, scalability, and scientific discovery effectiveness. Experimental analysis demonstrates that the proposed neuro-symbolic world model significantly improves reasoning transparency, reduces knowledge conflicts, enhances interdisciplinary knowledge integration, and accelerates scientific discovery compared with conventional deep learning and symbolic AI approaches. Furthermore, ethical governance mechanisms and uncertainty-aware reasoning ensure trustworthy deployment in critical research environments. The proposed framework establishes a scalable foundation for next-generation explainable artificial intelligence capable of autonomous scientific reasoning, continuous knowledge evolution, and reliable decision support across diverse scientific disciplines.

Keywords: Neuro-Symbolic AI, World Models, Explainable Artificial Intelligence, Scientific Discovery, Autonomous Knowledge Reasoning, Knowledge Graphs, Causal Inference, Logical Reasoning, Foundation Models, Cognitive Intelligence

Spatial Intelligence Framework for Real-Time Human-AI Collaboration in Physical Environments

Dr. Harish Kumar¹, John Peterson², Ananya Rao³

¹Professor, Department of AI & Data Science, PSG College of Technology, Coimbatore, India

²Principal Data Scientist, IBM Research, Yorktown Heights, New York, USA

³Assistant Professor, Department of CSE, Amrita Vishwa Vidyapeetham, Coimbatore, India

Abstract: The growing adoption of artificial intelligence in physical environments has created the need for intelligent systems capable of understanding three-dimensional spaces, interpreting human activities, and making context-aware decisions in real time. Traditional AI systems primarily process isolated sensory inputs and often lack comprehensive spatial awareness, limiting their effectiveness in dynamic environments where continuous human interaction is essential. This research proposes a Spatial Intelligence Framework (SIF) that enables real-time human-AI collaboration through integrated spatial perception, environmental reasoning, multimodal sensing, and adaptive decision-making. The proposed architecture combines computer vision, three-dimensional scene reconstruction, simultaneous localization and mapping (SLAM), multimodal foundation models, sensor fusion, and contextual reasoning to develop a continuously evolving digital representation of the surrounding environment. A spatial reasoning engine analyzes object relationships, human intentions, movement patterns, and environmental constraints to support safe, efficient, and collaborative interactions between humans and intelligent systems. The framework further incorporates reinforcement learning, semantic mapping, and continual learning mechanisms to improve adaptability under changing operational conditions. The proposed model is evaluated across multiple application domains, including autonomous robotics, smart manufacturing, intelligent healthcare facilities, warehouse automation, precision agriculture, disaster response, and smart infrastructure management. Performance evaluation focuses on spatial perception accuracy, collaboration efficiency, task completion rate, environmental understanding, computational latency, energy efficiency, human safety, and system reliability. Experimental results demonstrate that the proposed framework significantly improves spatial awareness, contextual understanding, collaborative decision-making, and autonomous task execution compared with conventional perception-based AI models. Additionally, explainable decision mechanisms and trust-aware interaction models enhance transparency and user confidence during human-AI collaboration. The proposed Spatial Intelligence Framework represents a significant advancement toward next-generation intelligent systems capable of operating safely, efficiently, and autonomously within complex real-world environments while fostering seamless collaboration between humans and artificial intelligence.

Keywords: Spatial Intelligence, Human-AI Collaboration, Multimodal AI, Computer Vision, Scene Understanding, Sensor Fusion, Simultaneous Localization and Mapping, Reinforcement Learning, Autonomous Systems, Intelligent Environments

Universal AI Operating Systems for Coordinating Autonomous Multi-Agent Ecosystems

Dr. Sophia Carter¹, Vikram Singh²

¹Associate Professor, Department of Computer Engineering, Carnegie Mellon University, Pittsburgh, Pennsylvania, USA

²Research Scholar, Department of Information Technology, Anna University, Chennai, India

Abstract: The rapid evolution of autonomous artificial intelligence has transformed isolated intelligent applications into interconnected ecosystems composed of multiple collaborative agents. However, existing AI infrastructures are primarily designed for individual models and lack a unified platform capable of coordinating heterogeneous autonomous agents operating across dynamic environments. This research proposes a **Universal AI Operating System (UAIOS)** that serves as an intelligent coordination layer for managing autonomous multi-agent ecosystems through adaptive resource allocation, distributed reasoning, collaborative planning, and real-time orchestration. The proposed architecture integrates agent lifecycle management, multimodal communication protocols, shared memory repositories, dynamic scheduling, semantic knowledge graphs, and policy-driven governance mechanisms to enable seamless cooperation among independent AI agents. Each agent specializes in perception, reasoning, planning, learning, or execution while the operating system continuously optimizes collaboration based on environmental context, computational availability, and task priorities. The framework incorporates reinforcement learning, decentralized consensus algorithms, continual learning, and trust-aware coordination strategies to improve system scalability, fault tolerance, and decision reliability. Performance evaluation is conducted across Industry 6.0 manufacturing, intelligent transportation, healthcare ecosystems, smart energy management, autonomous logistics, and large-scale enterprise automation. Evaluation metrics include coordination efficiency, task allocation accuracy, communication latency, computational resource utilization, fault recovery time, scalability, energy consumption, security resilience, and overall system performance. Experimental analysis demonstrates that the proposed Universal AI Operating System significantly enhances autonomous collaboration, minimizes operational conflicts, improves resource optimization, and supports efficient execution of complex distributed tasks compared with conventional multi-agent management frameworks. Furthermore, integrated explainability, security governance, and ethical policy enforcement ensure transparent and trustworthy autonomous operation. The proposed framework establishes a scalable foundation for future AI-native computing environments in which autonomous agents cooperate intelligently, continuously learn from shared experiences, and collectively solve complex real-world challenges across diverse application domains.

Keywords: Universal AI Operating System, Multi-Agent Systems, Autonomous Agents, Agent Orchestration, Distributed Artificial Intelligence, Resource Optimization, Reinforcement Learning, Continual Learning, Intelligent Automation, AI-Native Computing

Hybrid Quantum-Agentic Computing Architecture for High-Performance Scientific Problem Solving

Dr. Naveen Babu¹, Keerthana Raj²

¹Associate Professor, Department of Computer Science, Indian Institute of Technology Madras, Chennai, India

²AI Solutions Engineer, Zoho Corporation, Chennai, India

Abstract: The increasing complexity of scientific research requires computational frameworks capable of solving large-scale optimization, simulation, and reasoning problems that exceed the capabilities of conventional computing systems. While quantum computing offers exceptional processing potential for complex mathematical computations, it lacks autonomous reasoning and adaptive decision-making. Conversely, agentic artificial intelligence excels in planning, reasoning, and autonomous execution but remains constrained by classical computational resources for highly complex scientific tasks. This research proposes a Hybrid Quantum-Agentic Computing Architecture (HQACA) that integrates quantum computing with autonomous AI agents to create an intelligent computational ecosystem for next-generation scientific problem solving. The proposed framework combines quantum processors, large reasoning models, autonomous planning agents, knowledge graphs, quantum optimization algorithms, and continual learning mechanisms to distribute computational tasks based on their complexity and execution requirements. Agentic AI dynamically decomposes scientific problems, assigns optimization-intensive computations to quantum resources, performs classical reasoning, validates intermediate results, and continuously refines execution strategies through adaptive feedback. The framework further incorporates explainable reasoning, uncertainty estimation, resource scheduling, and trust-aware orchestration to ensure reliable and transparent scientific computing. The proposed architecture is evaluated across computational chemistry, drug discovery, climate modeling, materials science, astrophysics, genomic analysis, financial optimization, and complex engineering simulations. Performance evaluation includes computational efficiency, optimization accuracy, execution latency, scalability, resource utilization, reasoning consistency, solution quality, energy efficiency, and fault tolerance. Experimental analysis demonstrates that the proposed hybrid architecture significantly accelerates scientific computation, improves optimization performance, enhances autonomous reasoning, and reduces execution complexity compared with standalone quantum or classical AI systems. Furthermore, integrated security, ethical governance, and explainability mechanisms support responsible deployment in mission-critical scientific environments. The proposed Hybrid Quantum-Agentic Computing Architecture provides a scalable pathway toward future intelligent computing ecosystems capable of autonomously solving multidisciplinary scientific challenges while advancing the convergence of quantum computing and next-generation artificial intelligence.

Keywords: Hybrid Quantum Computing, Agentic Artificial Intelligence, Quantum Optimization, Scientific Computing, Large Reasoning Models, Autonomous Agents, Knowledge Graphs, Explainable AI, High-Performance Computing, Intelligent Problem Solving

Self-Verifying Foundation Models for Reliable Autonomous Research and Decision Making

Dr. William Harris¹, Riya Patel²

¹Professor, Department of Electrical Engineering, University of California, Berkeley, California, USA

²Research Associate, Department of Artificial Intelligence, IIT Hyderabad, India

Abstract: The increasing adoption of foundation models in scientific research and intelligent decision-making has demonstrated remarkable capabilities in knowledge generation, reasoning, and problem solving. However, these models frequently produce inaccurate information, hallucinated outputs, and inconsistent reasoning that reduce their reliability in high-stakes applications. This research proposes a Self-Verifying Foundation Model (**SVFM**) that enhances the trustworthiness of autonomous research and decision-making through continuous self-evaluation, evidence validation, logical consistency checking, and adaptive reasoning refinement. The proposed architecture integrates multimodal foundation models, knowledge graphs, retrieval-augmented reasoning, symbolic verification, uncertainty estimation, and confidence-aware decision modules to verify every generated output before final recommendation. During inference, the framework automatically compares generated responses with trusted knowledge sources, identifies reasoning conflicts, evaluates factual consistency, and performs recursive self-correction to minimize misinformation. A hierarchical validation engine continuously improves reasoning quality through feedback learning, enabling the model to adapt to evolving scientific knowledge while preserving previously validated information. The proposed framework is evaluated across scientific literature analysis, healthcare diagnosis, engineering design, financial forecasting, legal decision support, cybersecurity intelligence, and autonomous research assistance. Performance evaluation includes factual accuracy, reasoning consistency, hallucination reduction rate, verification efficiency, computational overhead, explainability score, trustworthiness, and autonomous decision quality. Experimental analysis demonstrates that the proposed self-verifying architecture significantly reduces incorrect predictions, improves logical reasoning, enhances evidence-based decision-making, and increases user confidence compared with conventional foundation models. Furthermore, the integration of transparent verification workflows, ethical governance policies, and uncertainty-aware reasoning supports responsible deployment in critical domains requiring high reliability. The proposed Self-Verifying Foundation Model establishes a robust foundation for future autonomous intelligence capable of producing accurate, explainable, and trustworthy decisions across multidisciplinary scientific and industrial applications.

Keywords: Self-Verifying Foundation Models, Autonomous Decision Making, Knowledge Verification, Explainable Artificial Intelligence, Retrieval-Augmented Generation, Knowledge Graphs, Scientific Research, Uncertainty Estimation, Trustworthy AI, Autonomous Reasoning

Cognitive Digital Organisms: Self-Adaptive Artificial Intelligence for Autonomous Computing Networks

Dr. Suresh Narayanan¹, Jennifer Brown², Dr. Daniel Lee³

¹Professor, Department of Electronics and Communication Engineering, National Institute of Technology Tiruchirappalli, India

²Associate Professor, Department of Computer Science, University of Texas at Austin, Texas, USA

³Senior AI Scientist, NVIDIA Corporation, Santa Clara, California, USA

Abstract: Future intelligent computing environments require autonomous systems capable of continuous adaptation, self-organization, and resilient decision-making without extensive human intervention. Traditional artificial intelligence architectures remain largely application-specific and lack the capability to autonomously evolve in response to dynamic operational conditions. This research introduces the concept of Cognitive Digital Organisms (CDOs), a novel self-adaptive artificial intelligence framework inspired by biological evolution and cognitive intelligence for autonomous computing networks. The proposed architecture combines autonomous agents, continual learning, distributed reasoning, evolutionary optimization, neural memory, and self-organizing network intelligence to create digital entities capable of sensing environmental changes, learning from experience, and autonomously modifying their behavior to improve long-term performance. Each digital organism independently monitors system conditions, exchanges knowledge with neighboring organisms through decentralized communication, and collectively optimizes network performance using cooperative intelligence. A cognitive adaptation engine integrates reinforcement learning, meta-learning, and causal reasoning to enable intelligent decision-making while maintaining system stability and operational resilience. The framework is evaluated across cloud computing, edge intelligence, smart manufacturing, autonomous transportation, cybersecurity defense, and intelligent communication networks. Performance metrics include adaptability, fault tolerance, resource optimization, learning efficiency, network resilience, computational scalability, decision accuracy, energy consumption, and autonomous recovery capability. Experimental results demonstrate that Cognitive Digital Organisms significantly improve network intelligence, operational reliability, collaborative learning, and self-healing performance compared with conventional distributed AI systems. The framework also incorporates explainable reasoning, security governance, and ethical policy enforcement to ensure transparent and trustworthy autonomous operation. This research establishes a biologically inspired paradigm for future autonomous computing ecosystems capable of continuous evolution, intelligent collaboration, and sustainable adaptation across highly dynamic digital infrastructures.

Keywords: Cognitive Digital Organisms, Self-Adaptive Artificial Intelligence, Autonomous Computing Networks, Continual Learning, Distributed Intelligence, Evolutionary Computing, Reinforcement Learning, Self-Organizing Systems, Intelligent Networks, Cognitive Computing

Embodied Foundation Models for Intelligent Human-Robot Collaboration in Industry 6.0

Dr. Rohit Verma¹, Olivia Martinez²

¹Associate Professor, Department of Artificial Intelligence, Delhi Technological University, New Delhi, India

²Research Scientist, Amazon Web Services, Seattle, Washington, USA

Abstract: The transition toward Industry 6.0 demands intelligent robotic systems capable of understanding human intentions, interpreting complex physical environments, and collaborating safely with human workers in real time. Conventional robotic control systems primarily rely on predefined programming and limited perception, restricting their ability to adapt to dynamic industrial environments. This research proposes an Embodied Foundation Model (EFM) that integrates multimodal perception, physical reasoning, language understanding, spatial intelligence, and autonomous decision-making to enable advanced human-robot collaboration in next-generation industrial ecosystems. The proposed architecture combines visual-language foundation models, tactile sensing, environmental mapping, large reasoning models, and continual learning mechanisms to develop a comprehensive understanding of physical workspaces and human activities. A collaborative reasoning engine continuously interprets worker intentions, predicts task sequences, evaluates environmental risks, and coordinates robotic actions through adaptive planning and real-time feedback. Reinforcement learning and imitation learning further enhance task execution by enabling robots to learn from both human demonstrations and autonomous experiences. The framework is evaluated in smart manufacturing, warehouse automation, precision assembly, logistics operations, healthcare robotics, and intelligent service environments. Performance evaluation includes collaboration efficiency, task completion accuracy, motion planning effectiveness, safety compliance, response latency, adaptability, energy efficiency, human satisfaction, and operational reliability. Experimental analysis demonstrates that the proposed embodied foundation model significantly improves collaborative productivity, environmental awareness, decision transparency, and autonomous task execution compared with conventional robotic intelligence systems. Furthermore, explainable reasoning mechanisms, trust-aware interaction models, and ethical governance policies ensure safe deployment within human-centered industrial environments. The proposed framework contributes to the realization of Industry 6.0 by establishing intelligent robotic systems capable of adaptive learning, natural collaboration, and autonomous decision-making while maintaining high standards of safety, efficiency, and operational trust.

Keywords: Embodied Foundation Models, Human-Robot Collaboration, Industry 6.0, Multimodal Artificial Intelligence, Spatial Intelligence, Reinforcement Learning, Intelligent Robotics, Autonomous Decision Making, Smart Manufacturing, Explainable Artificial Intelligence

Synthetic Universe Generation for Training Autonomous Physical Artificial Intelligence Systems

Dr. Akash Reddy¹, Sophia Wilson²

¹Professor, Department of Information Technology, VIT Chennai, India

²Assistant Professor, Department of Computer Science, Arizona State University, Arizona, USA

Abstract: The development of autonomous physical artificial intelligence requires extensive real-world data for training, validation, and deployment. However, collecting large-scale, diverse, and safely annotated datasets from physical environments is expensive, time-consuming, and often impractical for rare or hazardous scenarios. This research proposes a Synthetic Universe Generation (SUG) framework that creates highly realistic virtual environments for training autonomous physical AI systems using generative world models, physics-based simulations, and multimodal synthetic data generation. The proposed framework integrates diffusion models, neural rendering, digital twins, procedural environment generation, reinforcement learning, and physics engines to simulate complex real-world scenarios with high fidelity. Autonomous AI agents interact within these synthetic universes to learn navigation, reasoning, manipulation, collaboration, and adaptive decision-making without the risks associated with physical deployment. A dynamic environment engine continuously introduces unpredictable events, environmental variations, and object interactions to improve generalization and robustness. The framework is evaluated using autonomous robotics, self-driving vehicles, disaster response systems, industrial automation, healthcare robotics, and space exploration applications. Performance metrics include simulation realism, transfer learning efficiency, task completion accuracy, environmental adaptability, computational efficiency, safety validation, and autonomous learning performance. Experimental analysis demonstrates that synthetic universe training significantly reduces development costs, improves model robustness, accelerates deployment readiness, and enhances decision-making under uncertain conditions compared with conventional real-world-only training methods. Furthermore, explainable simulation analytics and ethical safety validation support responsible AI development. The proposed Synthetic Universe Generation framework provides a scalable and cost-effective foundation for training future autonomous physical intelligence capable of operating reliably across complex real-world environments.

Keywords: Synthetic Universe Generation, Physical Artificial Intelligence, Generative World Models, Digital Twins, Diffusion Models, Reinforcement Learning, Autonomous Robotics, Simulation Intelligence, Multimodal Learning, Autonomous Systems

Hierarchical Multi-Agent Intelligence Framework for Autonomous Enterprise Workflow Optimization

Dr. Benjamin Clark¹, Nithya Lakshmi²

¹Professor, School of Computing, Georgia Institute of Technology, Atlanta, Georgia, USA

²Research Scholar, Department of AI, Bharathiar University, Coimbatore, India

Abstract: Modern enterprises increasingly rely on artificial intelligence to automate complex business operations, yet existing workflow automation systems often function as isolated components with limited adaptability and coordination. This research proposes a Hierarchical Multi-Agent Intelligence Framework (HMAIF) that enables autonomous enterprise workflow optimization through cooperative agent collaboration, distributed reasoning, adaptive planning, and intelligent resource management. The proposed architecture organizes autonomous AI agents into hierarchical layers responsible for strategic planning, operational coordination, task execution, monitoring, and continuous optimization. Each intelligent agent specializes in business functions such as scheduling, financial analysis, customer service, cybersecurity, supply chain management, and human resource operations while communicating through a shared semantic knowledge repository. Reinforcement learning, large reasoning models, causal inference, and continual learning enable agents to dynamically allocate resources, predict workflow bottlenecks, resolve conflicts, and optimize enterprise performance under changing business conditions. The framework is evaluated across manufacturing, healthcare administration, banking, logistics, retail management, and cloud-based enterprise platforms. Performance evaluation includes workflow efficiency, resource utilization, decision accuracy, execution latency, operational scalability, collaboration effectiveness, fault tolerance, and business productivity. Experimental analysis demonstrates that the proposed hierarchical multi-agent architecture significantly improves workflow coordination, reduces operational delays, enhances resource optimization, and increases organizational adaptability compared with conventional robotic process automation and workflow management systems. Additionally, explainable reasoning mechanisms, trust-aware collaboration, and governance policies ensure transparent and responsible enterprise automation. The proposed framework establishes an intelligent foundation for AI-native organizations capable of autonomous workflow management, continuous adaptation, and sustainable digital transformation in Industry 6.0 environments.

Keywords: Multi-Agent Intelligence, Enterprise Automation, Workflow Optimization, Autonomous Agents, Distributed Artificial Intelligence, Reinforcement Learning, Resource Management, Intelligent Decision Making, Industry 6.0, Business Process Automation

Adaptive Neural Memory Networks for Lifelong Continual Learning in Foundation Models

Dr. Karthik Srinivasan¹, Grace Thomas², Aravind Raj³

¹Associate Professor, Department of Computer Science, IIT Kanpur, India

²AI Research Scientist, Intel Corporation, Santa Clara, California, USA

³Research Scholar, Department of AI, IIT Madras, Chennai, India

Abstract: Foundation models have achieved remarkable success in language understanding, reasoning, and multimodal intelligence; however, they continue to experience catastrophic forgetting when learning new information over extended periods. This limitation restricts their ability to function as continuously evolving intelligent systems. This research proposes an Adaptive Neural Memory Network (ANMN) that enables lifelong continual learning through hierarchical memory organization, dynamic knowledge consolidation, and adaptive retrieval mechanisms. The proposed framework integrates working memory, episodic memory, semantic memory, neural attention mechanisms, and external knowledge repositories to preserve previously acquired knowledge while efficiently incorporating new information. A memory optimization module continuously evaluates information relevance, compresses redundant representations, and strengthens important knowledge through reinforcement learning and self-supervised adaptation. Large reasoning models utilize the adaptive memory architecture to improve contextual understanding, long-term reasoning, and autonomous decision-making across evolving environments. The proposed framework is evaluated in scientific research assistance, healthcare diagnosis, financial intelligence, autonomous robotics, cybersecurity analytics, and educational technologies. Performance evaluation considers memory retention, continual learning efficiency, reasoning consistency, adaptation speed, computational overhead, scalability, contextual accuracy, and knowledge retrieval performance. Experimental analysis demonstrates that the proposed adaptive neural memory architecture significantly improves lifelong learning capabilities, minimizes catastrophic forgetting, enhances reasoning performance, and supports continuous knowledge evolution compared with conventional transformer-based foundation models. Furthermore, explainable memory visualization and trust-aware knowledge management enhance transparency and reliability in long-term AI deployment. This research contributes a scalable cognitive memory architecture capable of supporting future intelligent systems that continuously learn, reason, and adapt throughout their operational lifetime.

Keywords: Neural Memory Networks, Continual Learning, Foundation Models, Lifelong Learning, Knowledge Retention, Adaptive Artificial Intelligence, Reinforcement Learning, Cognitive Memory, Large Reasoning Models, Autonomous Intelligence

Trust-Calibrated Agentic Artificial Intelligence for Human-Centered Decision Support Systems

Dr. Rachel Moore¹, Pranav Nair²

¹Professor, Department of Computer Science, University of Illinois Urbana-Champaign, Illinois, USA

²Assistant Professor, Department of Artificial Intelligence, SRM University, Chennai, India

Abstract: As autonomous artificial intelligence becomes increasingly integrated into healthcare, finance, public administration, manufacturing, and critical infrastructure, establishing appropriate levels of human trust has become essential for safe and effective decision-making. Excessive trust may result in overreliance on AI recommendations, while insufficient trust can reduce the benefits of intelligent automation. This research proposes a Trust-Calibrated Agentic Artificial Intelligence (TCAAI) framework that dynamically measures, evaluates, and adjusts trust between human users and autonomous AI agents. The proposed architecture integrates explainable reasoning, uncertainty estimation, confidence scoring, behavioral analytics, causal inference, and adaptive interaction models to generate transparent recommendations aligned with human expectations and organizational policies. Autonomous agents continuously monitor decision quality, assess contextual risk, provide interpretable explanations, and adjust communication strategies according to user expertise and situational complexity. Reinforcement learning and continual feedback mechanisms enable the framework to improve collaboration quality while preventing automation bias and decision inconsistency. The proposed model is evaluated across clinical decision support, financial advisory systems, intelligent manufacturing, cybersecurity operations, emergency management, and public service applications. Performance evaluation includes trust calibration accuracy, decision reliability, user satisfaction, explainability, adaptation efficiency, collaboration effectiveness, computational performance, and ethical compliance. Experimental analysis demonstrates that the proposed framework significantly enhances human confidence, improves collaborative decision quality, reduces inappropriate AI reliance, and strengthens transparency compared with conventional decision support systems. Furthermore, integrated governance policies and responsible AI mechanisms ensure secure, fair, and accountable deployment. The proposed Trust-Calibrated Agentic AI framework provides a robust foundation for future human-centered intelligent systems that support reliable collaboration while maintaining transparency, adaptability, and ethical responsibility.

Keywords: Trust-Calibrated AI, Agentic Artificial Intelligence, Human-Centered AI, Decision Support Systems, Explainable AI, Uncertainty Estimation, Human-AI Collaboration, Responsible AI, Reinforcement Learning, Autonomous Decision Making

Physics-Grounded World Models for Autonomous Navigation in Unknown Dynamic Environments

Dr. Lakshmi Prasad¹, Jacob Miller²

¹Professor, Department of Computer Applications, Anna University, Chennai, India

²Machine Learning Engineer, Meta Platforms Inc., Menlo Park, California, USA

Abstract: Autonomous navigation in unknown and continuously changing environments remains one of the most challenging problems in artificial intelligence and robotics. Conventional navigation systems primarily depend on static maps, predefined rules, or large amounts of labeled data, limiting their adaptability in unpredictable situations. This research proposes a Physics-Grounded World Model (PGWM) that enables intelligent autonomous navigation through the integration of physical laws, multimodal perception, causal reasoning, and predictive environmental modeling. The proposed framework combines computer vision, LiDAR sensing, inertial measurements, physics simulation, spatial reasoning, and large reasoning models to construct an evolving representation of dynamic environments. A predictive planning engine estimates future environmental changes, identifies potential obstacles, evaluates alternative navigation strategies, and continuously updates decision policies using reinforcement learning and self-supervised adaptation. The framework further incorporates uncertainty estimation, semantic mapping, and continual learning to improve navigation accuracy in unfamiliar environments without requiring extensive retraining. The proposed model is evaluated using autonomous vehicles, intelligent drones, mobile service robots, industrial automation systems, disaster response platforms, and planetary exploration missions. Performance evaluation includes navigation accuracy, path optimization, collision avoidance, computational efficiency, environmental adaptability, energy consumption, decision latency, robustness, and safety compliance. Experimental results demonstrate that the proposed Physics-Grounded World Model significantly improves environmental understanding, long-term navigation planning, autonomous adaptability, and operational reliability compared with conventional deep learning-based navigation systems. Additionally, explainable decision mechanisms and trust-aware planning enhance transparency and user confidence in safety-critical applications. The proposed framework establishes a scalable foundation for future autonomous systems capable of intelligent navigation across highly dynamic and previously unseen physical environments.

Keywords: Physics-Grounded World Models, Autonomous Navigation, Spatial Intelligence, Reinforcement Learning, Computer Vision, LiDAR, Causal Reasoning, Dynamic Environments, Intelligent Robotics, Autonomous Systems

Energy-Adaptive Green Foundation Models for Sustainable Artificial Intelligence Infrastructure

Dr. David Wilson¹, Meena Krishnan², Dr. Vivek Sharma³

¹Associate Professor, Department of Computer Science, Purdue University, Indiana, USA

²Research Scholar, Department of CSE, IIT Delhi, India

³Professor, Department of Information Technology, PSG College of Technology, Coimbatore, India

Abstract: The rapid expansion of foundation models has significantly increased computational demands, resulting in substantial energy consumption, carbon emissions, and operational costs. As artificial intelligence becomes an essential component of modern digital infrastructure, developing environmentally sustainable AI systems has emerged as a critical research challenge. This study proposes an Energy-Adaptive Green Foundation Model (EAGF) framework that dynamically optimizes computational resources while maintaining high model performance across diverse application domains. The proposed architecture integrates energy-aware neural optimization, adaptive model scaling, dynamic token routing, intelligent workload scheduling, and carbon-aware resource allocation to minimize environmental impact without compromising predictive accuracy. A reinforcement learning controller continuously monitors hardware utilization, workload characteristics, renewable energy availability, and thermal conditions to optimize model execution in real time. The framework also incorporates model compression, sparse computation, knowledge distillation, and distributed edge-cloud collaboration to further reduce energy requirements. Experimental evaluation is conducted across large language processing, computer vision, autonomous robotics, scientific computing, healthcare analytics, and smart city applications. Performance metrics include energy consumption, carbon footprint, computational efficiency, model accuracy, inference latency, hardware utilization, operational cost, and sustainability index. Experimental results demonstrate that the proposed framework significantly reduces power consumption and greenhouse gas emissions while preserving competitive AI performance compared with conventional foundation model architectures. Additionally, transparent energy reporting and sustainability governance mechanisms enable responsible AI deployment aligned with global environmental objectives. The proposed Energy-Adaptive Green Foundation Model establishes a scalable foundation for developing environmentally responsible artificial intelligence infrastructure capable of supporting future intelligent computing ecosystems.

Keywords: Green Artificial Intelligence, Energy-Adaptive Computing, Foundation Models, Sustainable AI, Carbon-Aware Computing, Model Compression, Knowledge Distillation, Reinforcement Learning, Edge-Cloud Computing, Energy Optimization

Test-Time Knowledge Editing Framework for Continually Learning Large Language Models

Dr. Arun Prakash¹, Jessica Taylor²

¹Associate Professor, Department of Computer Science, IIT Roorkee, India

²Senior AI Engineer, Oracle Corporation, Austin, Texas, USA

Abstract: Large Language Models (LLMs) have achieved exceptional performance in reasoning, language understanding, and knowledge generation. However, updating their knowledge typically requires expensive retraining, making them less suitable for rapidly evolving real-world environments. This research proposes a Test-Time Knowledge Editing Framework (TTKEF) that enables LLMs to update factual knowledge dynamically during inference without modifying core model parameters. The proposed framework integrates retrieval-augmented reasoning, adaptive memory modules, confidence-aware verification, semantic consistency analysis, and lightweight knowledge editing mechanisms to ensure accurate and reliable responses. During inference, the framework identifies outdated or conflicting information, retrieves relevant evidence from trusted external repositories, validates contextual consistency, and temporarily updates the model's reasoning process without disrupting previously learned knowledge. Reinforcement learning and continual feedback mechanisms further optimize editing efficiency while minimizing hallucinations and catastrophic forgetting. The framework is evaluated across scientific literature analysis, healthcare knowledge systems, legal document processing, financial intelligence, educational platforms, and enterprise knowledge management. Performance evaluation includes knowledge update accuracy, inference latency, reasoning consistency, hallucination reduction, memory efficiency, adaptability, scalability, and user trust. Experimental analysis demonstrates that the proposed framework significantly improves knowledge freshness, response reliability, and reasoning transparency compared with conventional fine-tuning and parameter-editing techniques. Moreover, explainable editing mechanisms and governance policies ensure secure and responsible deployment in high-stakes applications. The proposed Test-Time Knowledge Editing Framework provides a practical solution for developing continuously evolving large language models capable of adapting to rapidly changing knowledge ecosystems while maintaining high computational efficiency and decision reliability.

Keywords: Test-Time Learning, Knowledge Editing, Large Language Models, Continual Learning, Retrieval-Augmented Generation, Hallucination Reduction, Adaptive Memory, Explainable AI, Knowledge Management, Artificial Intelligence

Collective Swarm Intelligence Using Cooperative Agentic AI for Distributed Autonomous Systems

Dr. Christopher Evans¹, Aishwarya Menon²

¹Professor, Department of Engineering, University of Michigan, Ann Arbor, Michigan, USA

²Research Associate, Department of Artificial Intelligence, VIT Vellore, India

Abstract: The increasing complexity of distributed autonomous systems requires intelligent coordination mechanisms that enable multiple autonomous agents to collaborate efficiently while adapting to dynamic environments. Traditional centralized control architectures often suffer from scalability limitations, communication bottlenecks, and single points of failure. This research proposes a Collective Swarm Intelligence (CSI) framework that combines cooperative Agentic Artificial Intelligence with decentralized swarm intelligence principles to achieve resilient and scalable autonomous coordination. The proposed architecture integrates autonomous reasoning agents, distributed knowledge sharing, adaptive communication protocols, consensus algorithms, reinforcement learning, and self-organizing behavioral models. Each intelligent agent independently perceives environmental changes, performs localized reasoning, exchanges contextual knowledge with neighboring agents, and collectively contributes to global decision-making through decentralized cooperation. Dynamic swarm optimization algorithms continuously adjust collaboration strategies according to task complexity, environmental uncertainty, and resource availability. The framework is evaluated across autonomous drone fleets, intelligent transportation systems, industrial robotics, disaster response operations, environmental monitoring, and smart logistics networks. Performance evaluation includes coordination efficiency, communication overhead, fault tolerance, scalability, resource utilization, decision accuracy, adaptability, response time, and mission success rate. Experimental analysis demonstrates that the proposed swarm intelligence framework significantly enhances collaborative problem-solving, operational resilience, distributed decision-making, and autonomous adaptability compared with conventional centralized multi-agent systems. Furthermore, explainable collaboration mechanisms, trust-aware communication, and secure coordination protocols ensure transparent and reliable deployment in safety-critical environments. The proposed Collective Swarm Intelligence framework establishes a scalable paradigm for future distributed autonomous ecosystems capable of intelligent cooperation, continuous adaptation, and efficient decentralized decision-making.

Keywords: Swarm Intelligence, Agentic Artificial Intelligence, Distributed Autonomous Systems, Multi-Agent Collaboration, Reinforcement Learning, Decentralized Intelligence, Autonomous Agents, Intelligent Coordination, Consensus Algorithms, Distributed Computing

Cross-Reality Artificial Intelligence Integrating Physical, Digital, and Spatial Computing Ecosystems

Dr. Senthil Kumar¹, Ethan Roberts²

¹Professor, Department of Information Technology, Kongu Engineering College, Erode, India

²Research Scientist, Apple Inc., Cupertino, California, USA

Abstract: The convergence of physical environments, digital infrastructures, and spatial computing technologies is transforming the future of intelligent human-computer interaction. However, existing artificial intelligence systems often operate independently across these environments, limiting seamless interoperability and contextual understanding. This research proposes a Cross-Reality Artificial Intelligence (CRAI) framework that integrates physical, digital, augmented, virtual, and mixed reality ecosystems into a unified intelligent computing architecture. The proposed framework combines multimodal foundation models, spatial intelligence, digital twins, environmental sensing, semantic scene understanding, and adaptive reasoning to create a continuous representation of interconnected real-world and virtual environments. Autonomous AI agents dynamically synchronize information across multiple reality layers, enabling intelligent collaboration, predictive decision-making, and immersive user interaction. Reinforcement learning and continual adaptation improve system responsiveness under changing environmental conditions while preserving contextual consistency across physical and digital spaces. The framework is evaluated across smart healthcare, industrial training, remote collaboration, autonomous robotics, education, urban planning, and emergency response applications. Performance metrics include spatial understanding accuracy, synchronization efficiency, contextual consistency, interaction quality, computational performance, scalability, latency, user experience, and operational reliability. Experimental analysis demonstrates that the proposed Cross-Reality AI framework significantly enhances immersive intelligence, collaborative decision-making, environmental awareness, and real-time interaction compared with conventional isolated AI systems. Additionally, privacy-aware data management, explainable reasoning, and ethical governance mechanisms support trustworthy deployment across heterogeneous digital ecosystems. The proposed framework establishes a comprehensive foundation for future intelligent cross-reality platforms capable of seamlessly integrating physical and virtual worlds through adaptive artificial intelligence.

Keywords: Cross-Reality Artificial Intelligence, Spatial Computing, Digital Twins, Mixed Reality, Multimodal AI, Autonomous Agents, Spatial Intelligence, Human-Computer Interaction, Continual Learning, Intelligent Environments

Autonomous AI Scientists for Scientific Hypothesis Generation and Experimental Validation

Dr. Abigail Scott¹, Rahul Desai², Dr. Kavya Raman³

¹Associate Professor, Department of Computer Science, Cornell University, Ithaca, New York, USA

²Research Scholar, Department of AI, NIT Surathkal, India

³Professor, Department of Artificial Intelligence, Anna University, Chennai, India

Abstract: Scientific discovery increasingly relies on large-scale data analysis, interdisciplinary knowledge integration, and complex experimental design, creating challenges that exceed traditional research methodologies. This research proposes an Autonomous AI Scientist (AAIS) framework capable of independently generating scientific hypotheses, designing experiments, analyzing results, and refining knowledge through iterative reasoning. The proposed architecture integrates foundation models, large reasoning models, neuro-symbolic reasoning, causal inference, knowledge graphs, retrieval-augmented generation, and autonomous planning to support end-to-end scientific research workflows. The system continuously analyzes multidisciplinary literature, identifies unresolved research gaps, formulates testable hypotheses, designs optimized experimental procedures, predicts expected outcomes, and validates findings using both simulated and real-world evidence. Reinforcement learning and continual learning mechanisms enable the AI scientist to improve reasoning quality, experimental efficiency, and knowledge integration over time. The framework is evaluated across drug discovery, materials science, genomics, climate science, computational chemistry, biomedical research, and engineering design. Performance evaluation includes hypothesis novelty, experimental success rate, reasoning consistency, knowledge integration accuracy, computational efficiency, explainability, reproducibility, and scientific impact. Experimental analysis demonstrates that the proposed framework significantly accelerates scientific discovery, reduces research costs, enhances interdisciplinary collaboration, and improves evidence-based decision-making compared with conventional AI-assisted research tools. Furthermore, transparent reasoning, ethical governance, and human oversight mechanisms ensure responsible deployment in scientific environments. The proposed Autonomous AI Scientist framework provides a transformative foundation for next-generation intelligent research systems capable of continuously advancing scientific knowledge through autonomous reasoning, experimentation, and discovery.

Keywords: Autonomous AI Scientist, Scientific Discovery, Hypothesis Generation, Experimental Validation, Large Reasoning Models, Neuro-Symbolic AI, Knowledge Graphs, Causal Inference, Continual Learning, Intelligent Research Systems

Self-Healing Artificial Intelligence Infrastructure Using Autonomous Multi-Agent Orchestration

Dr. Mohan Raj¹, Sarah Thompson²

¹Associate Professor, Department of AI & ML, SASTRA University, Thanjavur, India

²AI Systems Engineer, Tesla Inc., Palo Alto, California, USA

Abstract: As artificial intelligence infrastructures become increasingly distributed across cloud, edge, and hybrid computing environments, maintaining system reliability, resilience, and continuous availability has become a significant challenge. Traditional infrastructure management relies heavily on manual monitoring and reactive maintenance, resulting in increased downtime and operational inefficiencies. This research proposes a Self-Healing Artificial Intelligence Infrastructure (SHAI) framework that employs autonomous multi-agent orchestration to detect, diagnose, and recover from infrastructure failures with minimal human intervention. The proposed architecture integrates intelligent monitoring agents, predictive analytics, causal reasoning, distributed knowledge graphs, reinforcement learning, and adaptive resource orchestration to enable proactive infrastructure management. Autonomous agents continuously observe hardware performance, network behavior, software services, and cybersecurity events while collaboratively identifying anomalies, predicting failures, and initiating corrective actions. Dynamic orchestration mechanisms optimize workload migration, service recovery, resource allocation, and fault isolation to maintain uninterrupted system performance. The framework is evaluated across cloud computing platforms, edge intelligence networks, smart manufacturing systems, healthcare infrastructures, financial data centers, and autonomous transportation ecosystems. Performance evaluation includes fault detection accuracy, recovery time, system availability, resource utilization, scalability, energy efficiency, operational resilience, and cybersecurity robustness. Experimental analysis demonstrates that the proposed self-healing infrastructure significantly reduces service interruptions, improves infrastructure reliability, enhances autonomous recovery capabilities, and minimizes maintenance costs compared with conventional infrastructure management approaches. Furthermore, explainable orchestration strategies, security governance, and ethical operational policies ensure transparent and trustworthy infrastructure management. The proposed framework establishes a scalable foundation for resilient AI-native computing environments capable of autonomous maintenance, intelligent adaptation, and continuous service availability.

Keywords: Self-Healing AI, Autonomous Infrastructure, Multi-Agent Orchestration, Cloud Computing, Edge Computing, Predictive Maintenance, Reinforcement Learning, Intelligent Monitoring, Fault Recovery, AI-Native Infrastructure

Foundation Models for Universal Scientific Simulation across Interdisciplinary Research Domains

Dr. Noah Mitchell¹, Pooja Kapoor²

¹Professor, Department of Computer Engineering, Duke University, North Carolina, USA

²Research Scholar, Department of CSE, IIT Guwahati, India

Abstract: Scientific simulation plays a fundamental role in advancing knowledge across disciplines including healthcare, climate science, engineering, physics, chemistry, and biology. However, existing simulation platforms are typically domain-specific and require extensive customization for each research application. This research proposes a Universal Scientific Simulation Framework (USSF) powered by foundation models capable of generating, adapting, and optimizing scientific simulations across multiple research domains through unified knowledge representation and autonomous reasoning. The proposed architecture integrates multimodal foundation models, physics-informed neural networks, scientific knowledge graphs, causal reasoning, reinforcement learning, and adaptive simulation engines to construct highly accurate virtual representations of complex physical and biological systems. The framework automatically interprets scientific objectives, selects appropriate computational models, predicts simulation outcomes, and continuously refines results through iterative learning and evidence validation. Cross-domain knowledge transfer further enables the system to leverage discoveries from one discipline to accelerate research in others. The proposed framework is evaluated across computational biology, climate prediction, materials science, fluid dynamics, pharmaceutical research, aerospace engineering, and renewable energy optimization. Performance evaluation includes simulation accuracy, computational efficiency, scalability, adaptability, knowledge transfer effectiveness, reasoning consistency, reproducibility, and scientific discovery potential. Experimental analysis demonstrates that the proposed foundation model significantly improves simulation fidelity, reduces computational costs, enhances interdisciplinary collaboration, and accelerates scientific innovation compared with conventional domain-specific simulation platforms. Additionally, explainable simulation workflows, uncertainty estimation, and ethical governance mechanisms ensure trustworthy deployment in critical scientific environments. The proposed Universal Scientific Simulation Framework establishes a transformative foundation for future AI-driven scientific research by enabling intelligent, scalable, and interdisciplinary simulation capabilities.

Keywords: Foundation Models, Scientific Simulation, Physics-Informed Neural Networks, Knowledge Graphs, Computational Science, Scientific Computing, Causal Reasoning, Reinforcement Learning, Interdisciplinary Research, Artificial Intelligence

Causal World Models for Autonomous Decision Intelligence in Complex Dynamic Systems

Dr. Gokul Krishna¹, Emily Davis², Arjun Patel³

¹Professor, Department of Information Technology, SRM Institute of Science and Technology, Chennai, India

²Senior Software Engineer, Cisco Systems, San Jose, California, USA

³Research Scholar, Department of AI, Amrita Vishwa Vidyapeetham, Coimbatore, India

Abstract: Autonomous intelligent systems operating in real-world environments must understand not only correlations but also the underlying causal relationships that govern dynamic processes. Conventional artificial intelligence models often rely on statistical learning techniques that struggle to generalize under changing conditions and unforeseen events. This research proposes a Causal World Model (CWM) framework that integrates causal inference, world modeling, multimodal perception, and adaptive reasoning to support reliable autonomous decision intelligence in complex dynamic systems. The proposed architecture combines knowledge graphs, structural causal models, reinforcement learning, probabilistic reasoning, and temporal memory to construct an evolving representation of environmental cause-and-effect relationships. The framework continuously predicts future system states, evaluates intervention strategies, identifies hidden causal dependencies, and adapts decision policies through continual learning. A reasoning engine generates transparent explanations that justify autonomous decisions based on interpretable causal pathways rather than statistical associations alone. The proposed framework is evaluated across autonomous transportation, healthcare diagnosis, industrial automation, financial risk analysis, climate forecasting, disaster management, and smart infrastructure applications. Performance evaluation includes causal inference accuracy, decision reliability, prediction consistency, adaptability, computational efficiency, explainability, robustness, and autonomous planning performance. Experimental analysis demonstrates that the proposed Causal World Model significantly improves decision quality, long-term prediction accuracy, environmental adaptability, and reasoning transparency compared with conventional deep learning models. Furthermore, integrated ethical governance, uncertainty estimation, and trust-aware reasoning mechanisms support responsible deployment in safety-critical environments. The proposed framework contributes a scalable foundation for future autonomous intelligence capable of understanding complex causal relationships and making robust, evidence-based decisions across diverse application domains.

Keywords: Causal World Models, Autonomous Decision Intelligence, Causal Inference, Reinforcement Learning, Knowledge Graphs, Explainable Artificial Intelligence, Dynamic Systems, Structural Causal Models, Continual Learning, Intelligent Decision Making

Xeromorphic Argentic Computing for Ultra-Efficient Next-Generation Artificial Intelligence

Dr. Robert Green¹, Harini Suresh²

¹Associate Professor, Department of Computer Science, University of Florida, Gainesville, Florida, USA

²Assistant Professor, Department of Artificial Intelligence, Sabetha Engineering College, Chennai, India

Abstract: The exponential growth of artificial intelligence has increased computational demands beyond the capabilities of conventional von Neumann computing architectures. Xeromorphic computing offers a biologically inspired alternative that mimics the structure and functionality of the human brain while significantly reducing energy consumption and processing latency. This research proposes a Xeromorphic Argentic Computing (NAC) framework that combines xeromorphic hardware with autonomous argentic intelligence to develop ultra-efficient artificial intelligence systems capable of adaptive reasoning, continual learning, and autonomous decision-making. The proposed architecture integrates spiking neural networks, event-driven computation, large reasoning models, adaptive memory modules, reinforcement learning, and decentralized autonomous agents to perform intelligent computation using energy-efficient neural processing. Argentic control modules coordinate perception, planning, reasoning, and execution while dynamically allocating computational resources based on environmental complexity and task priorities. The framework is evaluated across edge computing, autonomous robotics, intelligent healthcare devices, smart manufacturing, autonomous vehicles, environmental monitoring, and space exploration systems. Performance evaluation includes energy efficiency, computational latency, learning adaptability, reasoning accuracy, scalability, hardware utilization, fault tolerance, and operational reliability. Experimental analysis demonstrates that the proposed xeromorphic argentic framework significantly reduces power consumption while maintaining high levels of autonomous reasoning and adaptive intelligence compared with conventional GPU-based AI systems. Furthermore, explainable decision-making mechanisms, trust-aware learning strategies, and ethical governance policies ensure responsible deployment across mission-critical applications. The proposed Xeromorphic Argentic Computing framework provides a scalable pathway toward sustainable, intelligent, and energy-efficient computing infrastructures for future autonomous AI ecosystems.

Keywords: Xeromorphic Computing, Argentic Artificial Intelligence, Spiking Neural Networks, Edge Computing, Reinforcement Learning, Adaptive Memory, Autonomous Agents, Energy-Efficient AI, Intelligent Computing, Next-Generation AI

World Model-Based Autonomous Cyber Defense Using Intelligent Multi-Agent Threat Reasoning

Dr. Vinay Kumar¹, Laura White²

¹Professor, Department of AI and Robotics, IIT Hyderabad, India

²Research Scientist, Qualcomm Technologies, San Diego, California, USA

Abstract: The increasing sophistication of cyber threats requires defense systems capable of predicting, reasoning, and autonomously responding to rapidly evolving attack strategies. Conventional cybersecurity solutions primarily rely on signature-based detection and reactive analysis, making them ineffective against advanced persistent threats and zero-day attacks. This research proposes a World Model-Based Autonomous Cyber Defense (WMACD) framework that integrates world models, intelligent multi-agent reasoning, causal threat analysis, and adaptive defense mechanisms to enable proactive cybersecurity operations. The proposed architecture combines foundation models, knowledge graphs, attack simulation environments, reinforcement learning, threat intelligence repositories, and distributed autonomous agents to construct a continuously evolving representation of cyber ecosystems. Intelligent agents collaboratively monitor network activities, identify hidden attack patterns, predict adversarial behavior, evaluate potential risks, and autonomously coordinate defense strategies. A causal reasoning engine analyzes attack propagation pathways, estimates system vulnerabilities, and generates explainable security recommendations to support transparent decision-making. The framework is evaluated across cloud infrastructures, enterprise networks, Internet of Things ecosystems, industrial control systems, financial institutions, healthcare networks, and national critical infrastructure. Performance evaluation includes threat detection accuracy, attack prediction capability, response latency, false-positive rate, system resilience, scalability, computational efficiency, and autonomous recovery performance. Experimental analysis demonstrates that the proposed framework significantly enhances cyber resilience, proactive threat mitigation, autonomous incident response, and decision transparency compared with conventional AI-based security systems. Additionally, integrated ethical governance, privacy protection, and explainable security analytics ensure trustworthy deployment in mission-critical digital infrastructures. The proposed World Model-Based Autonomous Cyber Defense framework establishes a robust foundation for future intelligent cybersecurity ecosystems capable of adaptive reasoning, autonomous protection, and continuous threat evolution analysis.

Keywords: Autonomous Cyber Defense, World Models, Multi-Agent Systems, Cyber Threat Intelligence, Reinforcement Learning, Knowledge Graphs, Explainable Cybersecurity, Network Security, Intelligent Threat Reasoning, Artificial Intelligence

Large Reasoning Models for Explainable Scientific Computing Beyond Conventional Large Language Models

Dr. Andrew Collins¹, Deepika Iyer²

¹Professor, Department of Computer Science, University of Washington, Seattle, Washington, USA

²Research Scholar, Department of AI, Anna University, Chennai, India

Abstract: The rapid evolution of artificial intelligence has shifted research from conventional Large Language Models (LLMs) toward Large Reasoning Models (LRMs) that emphasize structured reasoning, logical inference, and transparent decision-making. Although LLMs demonstrate exceptional language generation capabilities, they frequently produce hallucinations, inconsistent reasoning, and limited interpretability, restricting their applicability in scientific computing and high-risk decision environments. This research proposes a Large Reasoning Model (LRM) architecture designed to support explainable scientific computing through hierarchical reasoning, causal inference, symbolic knowledge integration, and adaptive problem-solving. The proposed framework integrates multimodal scientific datasets, knowledge graphs, mathematical reasoning engines, retrieval-augmented reasoning, and continual learning mechanisms to generate accurate, verifiable, and transparent scientific outcomes. A multi-stage reasoning pipeline decomposes complex scientific problems into interpretable reasoning steps, validates intermediate conclusions using external evidence, and produces confidence-aware explanations for every computational decision. The framework is evaluated across computational biology, climate modeling, materials science, healthcare analytics, engineering optimization, financial forecasting, and space science. Performance evaluation includes reasoning accuracy, computational efficiency, explainability score, hallucination reduction, knowledge consistency, scalability, reproducibility, and scientific reliability. Experimental analysis demonstrates that the proposed LRM significantly improves logical consistency, scientific accuracy, evidence-based reasoning, and transparency compared with conventional LLM-based scientific computing systems. Furthermore, uncertainty estimation, ethical governance, and trust-aware validation mechanisms support responsible deployment in mission-critical scientific applications. The proposed Large Reasoning Model establishes a scalable foundation for next-generation explainable scientific computing capable of autonomous reasoning, interdisciplinary knowledge integration, and trustworthy decision intelligence beyond traditional language-based artificial intelligence.

Keywords: Large Reasoning Models, Explainable Scientific Computing, Scientific Artificial Intelligence, Causal Reasoning, Knowledge Graphs, Retrieval-Augmented Reasoning, Foundation Models, Continual Learning, Trustworthy AI, Autonomous Scientific Intelligence

AI-Native Autonomous Computing Platforms for Future Intelligent Digital Ecosystems

Dr. Praveen Natarajan¹, Kevin Brooks²

¹Associate Professor, Department of Computer Science and Engineering, IIT Indore, India

²Machine Learning Scientist, OpenAI, San Francisco, California, USA

Abstract: The future of digital transformation requires computing platforms where artificial intelligence is not merely an application but the fundamental operating principle governing infrastructure, resource management, security, and autonomous decision-making. Conventional computing platforms depend on predefined system administration and manual optimization, limiting their ability to adapt to dynamic workloads and evolving user requirements. This research proposes an AI-Native Autonomous Computing Platform (ANACP) that integrates foundation models, autonomous agents, intelligent orchestration, continual learning, and predictive resource optimization to create self-managing digital ecosystems. The proposed architecture combines distributed computing, cloud-edge collaboration, knowledge graphs, reinforcement learning, and adaptive infrastructure management to continuously monitor system performance, predict workload fluctuations, allocate computational resources, optimize energy consumption, and autonomously recover from failures. Intelligent service agents collaborate through decentralized reasoning while maintaining secure communication and transparent governance policies. The framework is evaluated across cloud computing infrastructures, smart cities, enterprise digital transformation, healthcare information systems, intelligent manufacturing, financial technology platforms, and large-scale scientific computing environments. Performance evaluation considers resource utilization, scalability, operational efficiency, fault recovery, energy optimization, cybersecurity resilience, computational latency, and autonomous management effectiveness. Experimental results demonstrate that the proposed AI-native platform significantly improves infrastructure intelligence, reduces operational complexity, enhances service reliability, and supports continuous adaptation compared with conventional cloud management architectures. Additionally, explainable infrastructure management, ethical AI governance, and trust-aware orchestration mechanisms ensure responsible deployment in critical computing environments. The proposed AI-Native Autonomous Computing Platform establishes a scalable architecture for future intelligent digital ecosystems capable of self-management, autonomous optimization, and resilient operation across highly distributed computational infrastructures.

Keywords: AI-Native Computing, Autonomous Computing Platforms, Intelligent Infrastructure, Distributed Artificial Intelligence, Reinforcement Learning, Cloud-Edge Computing, Autonomous Agents, Resource Optimization, Digital Ecosystems, Intelligent Orchestration

Human-Centered Artificial General Intelligence Framework for Safe and Responsible Autonomous Systems

Dr. Matthew Robinson¹, Shruthi Narayanan², Dr. Rajesh Kumar³

¹Professor, Department of Computer Engineering, University of Maryland, College Park, Maryland, USA

²Research Scholar, Department of Information Technology, Bharath Institute of Higher Education and Research, Chennai, India

³Professor, Department of Computer Science, IIT Kharagpur, India

Abstract: Artificial General Intelligence (AGI) represents the next frontier of intelligent computing, aiming to develop systems capable of learning, reasoning, and adapting across diverse tasks with human-level cognitive flexibility. Despite rapid progress in foundation models and autonomous agents, ensuring the safety, transparency, ethical alignment, and societal acceptance of AGI remains one of the most significant challenges in artificial intelligence research. This study proposes a Human-Centered Artificial General Intelligence (HC-AGI) framework that integrates cognitive reasoning, value alignment, explainable decision-making, continual learning, and adaptive governance to ensure that autonomous systems consistently operate in accordance with human objectives and ethical principles. The proposed architecture combines large reasoning models, neuro-symbolic intelligence, causal world models, multimodal perception, reinforcement learning, and human-in-the-loop feedback mechanisms to support reliable and interpretable decision-making across dynamic environments. A comprehensive governance layer continuously evaluates fairness, transparency, accountability, privacy protection, and risk mitigation while monitoring system behavior and adapting policies according to evolving societal and regulatory requirements. The framework is evaluated across healthcare, education, scientific research, intelligent transportation, industrial automation, public administration, and disaster response systems. Performance evaluation includes reasoning capability, alignment accuracy, explainability, adaptability, robustness, ethical compliance, computational efficiency, human trust, and operational safety. Experimental analysis demonstrates that the proposed Human-Centered AGI framework significantly improves collaborative intelligence, decision transparency, value alignment, and long-term system reliability compared with existing autonomous AI architectures. Furthermore, integrated responsible AI governance and continuous human oversight provide a scalable pathway toward the safe deployment of future general-purpose intelligent systems. The proposed framework contributes to the realization of trustworthy AGI that enhances human capabilities while ensuring ethical, secure, and socially beneficial autonomous intelligence.

Keywords: Artificial General Intelligence, Human-Centered AI, Responsible Artificial Intelligence, Large Reasoning Models, Neuro-Symbolic Intelligence, Causal World Models, Explainable AI, Human-in-the-Loop Learning, Ethical AI Governance, Autonomous Intelligent Systems