

CYBER TIMES[®]

ISSN: 2278-7518

INTERNATIONAL JOURNAL OF TECHNOLOGY AND MANAGEMENT

Volume 19 - Issue 01, October 2025 - March 2026
Bi-Annual Double Blind Peer Reviewed Refereed Journal



CYBER  TIMES[®]
(Leader in innovative Tech-World)

Cyber Times International Journal of Technology & Management

Volume 19 - Issue 01, October 2025 – March 2026
ISSN: 2278-7518

EDITOR-IN-CHIEF

Dr. Anup Girdhar

EDITORIAL ADVISORY BOARD

Dr. Sushila Madan
Dr. A.K. Saini
Advocate Mukul Girdhar
Ms. Sonia Girdhar

EXECUTIVE EDITORS

Ms. Samridhi Girdhar
Mr. Rakesh Laxman Patil



Scientific Journal Impact Factor Value for 2024 = 6.007

“Cyber Times International Journal of Technology & Management”. All rights reserved. No part of this journal may be reproduced, republished, stored, or transmitted in any form or by any means, electronic, mechanical, photocopying, recording, or otherwise, without the prior permission of the publisher in writing. Any person who does any unauthorized act in relation to this journal publication may be liable to criminal prosecution and civil claims for damages.

Editorial Office & Administrative Address:

Delhi:

The Editor,
A19/1, Mansa Ram Park,
New Delhi-110059.

ISSN: 2278-7518

Phone: +91-9811485729, +91-9312903095

Website: <https://journal.cybertimes.in>

Email: editor@cybertimes.in

Disclaimer: Views and information expressed in the Research Papers or Articles are those of the respective authors. **“Cyber Times International Journal of Technology & Management”**, its Editorial Board, Editor and Publisher (Cyber Times) disclaim the Responsibility and Liability for any statement of fact or opinion made by the contributors. The content of the papers are written by their respective authors. The originality and authenticity of the papers and the explanation of information and views expressed therein are the sole responsibility of the authors. However, effort is made to acknowledge source material relied upon or referred to, however; **“Cyber Times International Journal of Technology & Management”** does not accept any responsibility for any unintentional mistakes & errors.

From the Editor's Desk

At the outset, I take this opportunity to express my sincere gratitude to all the Editorial Board Members, Editors, Peer Review Members, contributors, and readers for making *Cyber Times International Journal of Technology & Management* an outstanding success. Their unwavering support, dedication, and commitment to academic excellence have significantly contributed to the growth and reputation of the journal.

We are pleased to present **Volume 19 – Issue 01** of *Cyber Times International Journal of Technology & Management*. This issue features a collection of high-quality research papers and scholarly articles that reflect contemporary developments, innovative ideas, and critical insights across emerging areas of Technology, Management, Law, Education, and other multidisciplinary domains. The diversity of topics covered in this issue highlights the increasing importance of interdisciplinary research in addressing global challenges and opportunities.

The overwhelming response received from researchers, authors, academicians, law-enforcement agencies, and industry professionals for submitting their research papers and articles is deeply appreciated and duly acknowledged across the globe. Their valuable contributions have enriched the journal's content and strengthened its role as a platform for disseminating knowledge, fostering innovation, and encouraging scholarly dialogue among academia, industry, and society.

On behalf of the Editorial Team, I extend my heartfelt thanks to all authors for their valuable research contributions and to our reviewers for their constructive evaluations that help maintain the highest standards of publication quality. We hope that the research published in this issue will inspire further inquiry, collaboration, and advancement in various fields of study, while continuing to serve as a meaningful resource for our readers worldwide.

We look forward to receive your valuable and future contributions to make this journal a joint endeavor.

With Warm Regards,



Dr. ANUP GIRDHAR

Editor-In-Chief

Cyber Times International Journal of Technology & Management

General Information

- ***“Cyber Times International Journal of Technology & Management”*** is published bi-annually. All editorial and administrative correspondence for publication should be addressed to The Editor, Cyber Times.
- The Abstracts received for the final publication are screened by the Evaluation Committee for approval and only the selected Papers/ Abstracts will be published in each edition. Further information is available in the **“Guidelines for paper Submission”** section.
- Annual Subscription details for obtaining the print copy of the journal are provided separately and the interested persons may avail the same accordingly after filling the Annual subscription form.
- This journal is meant for education, reference and learning purposes. The author(s) of this of the book has/have taken all reasonable care to ensure that the contents of the book do not violate any existing copyright or other intellectual property rights of any person/ company/ institution in any manner whatsoever. In the event the author(s) has/have been unable to track any source and if any copyright has been inadvertently infringed, please notify the publisher in writing for the corrective action.
- Copyright © ***“Cyber Times International Journal of Technology & Management”***. All rights reserved. No part of this journal may be reproduced, republished, stored, or transmitted in any form or by any means, electronic, mechanical, photocopying, recording, or otherwise, without the prior permission of the publisher in writing. Any person who does any unauthorized act in relation to this journal publication may be liable to criminal prosecution and civil claims for damages.
- **Other Publications:**
 - Cyber Times Newspaper (English) – RNI No: DELENG/2008/25470
 - Cyber Times Newspaper (Hindi) – RNI No. DELHIN/1999/00462
- **Printed & Published by:** Cyber Times
A19/1, Mansa Ram Park, New Delhi-110059

CTIJTM Editorial Advisory Board Members

Name	Designation, Organization/ University	Country
Dr. Sushila Madan	Associate Professor, Delhi University	India
Dr. A. K. Saini	Professor, GGS IP University	India
Mr. J. R. Ahuja	Former Consultant, AICTE	India
Mr. Mukul Girdhar	Advocate Delhi High Court	India
Mr. Geetesh Madan	Q.A. Consultant with Tesco Bank, Newcastle	UK
Dr. Deepak Shikarpur	Chairman Board of Studies, Pune University	India
Dr. B. B. Ahuja	Deputy Director, COE - Pune	India
Prof. M. N. Hoda	Director, Bharati Vidyapeeth's (BVICAM)	India
Dr. S. C. Gupta	Director, NIEC, GGS IP University	India
Dr. S. K. Gupta	Professor, IIT Delhi	India
Dr. K. S. Wani	Principal, SSBT's COET, Bambhori, Jalgaon	India
Dr. K. V. Arya	Associate Professor, IIITM, Gwalior	India
BRIG. Dr. S.S. Narula	Director, Gitarattan International Bussiness School	India
Dr. Sarika Sharma	Director, JSPM'S ENIAC Institute of CA, Pune	India
Dr. S.K.M. Bhagat	Prof. & Head, MIT Academy of Engg., Pune	India
Dr. Jack Ajowi	Jaramogi Oginga Odinga University of Sci. & Tech.	Kenya
Dr. Srinivas Sampalli	Professor, Dalhousie University, Halifax	Canada
Dr. Ijaz A. Qureshi	V.P. Academic Affairs, JFK Inst. of Tech. and Mgmt.	Pakistan
Aryya Bhattacharyya	Director, CIP, Columbus State University	US
Dr. M. M. Schiraldi	Assistant Professor, 'Tor Vergata' University of Rome	Italy

Executive Editorial Advisory Board Members

Name	Designation, Organization/ University	Country
Ms. Kanika Trehan	Editor - Cyber Times, New Delhi	India
Mr. Rakesh Laxman Patil	Editor - Cyber Times, Pune	India
Adv. Tushar Kale	Cyber Lawyer, Pune	India
Adv. Neeraj Aarora	Cyber Lawyer, New Delhi	India
Mr. Sanjeev Sehgal	HOD, SJP Polytech, Damla, Haryana	India
Mr. Rajinder Kumar Bajaj	GM, Satake India Engg. Pvt. Ltd., (Japan)	India
Dr. B. M. Patil	Associate Professor MIT, Pune	India
Dr. Rajesh S. Prasad	Professor, DCOER, Pune University	India
Dr. Binod Kumar	Associate Professor, MIT Academy of Engg, Pune	India
Prof. Dr. M. Husain	HOD, SSBT's COET, Bambhori, Jalgaon	India
Prof. Dr. U. S. Bhadade	HOD, SSBT's COET, Bambhori, Jalgaon	India
Dr. V. N. Wadekar	Prof. & Head, MIT college of Engg. CMSR, Pune	India
Dr. M.D. Goudar	Associate Prof. & Head, Pune University	India
Dr. Mohd. Rizwan Alam	Sr. Lecturer, Amity University	Dubai
Prof. Jagannath Aghav	Professor & Head, CSE & IT, COE - Pune	India

Disclaimer: The names, affiliations, and designations of the Editorial Board Members published in this journal are based on the information provided at the time of their association with *Cyber Times International Journal of Technology & Management*. Members may subsequently change their positions, affiliations, or professional designations. The journal does not guarantee the continued accuracy of such details after publication.

Cyber Times International Journal of Technology and Management - CTIJTM

Volume 19 - Issue 01

CONTENTS

S.No.	Title	Page No.
1.	The Changing Landscape of Airline Retailing and the myths surrounding customer preferences <i>Dr. C. Sunanda Yadav</i>	01
2.	Computer Science Skills: Analyzing the Gap Between Academia and Industry <i>Asst. Prof. Sudhir Suman & Dr. Anup Girdhar</i>	08
3.	A Doctrinal Analysis of the Impact of Digitalisation on Legal Education in India: Legal Framework, Reforms, and Challenges <i>Mrs. Sapana Jaiswal</i>	14
4.	Artificial Intelligence and Legal Personhood: A Study of Socio-Legal Implications in India <i>Asst. Prof. Seema A. Patil</i>	21
5.	Environmental Degradation and Management: An Analytical Study of Causes, Impacts, and Mitigation Strategies <i>Dr. Kalpana Ghatpande & Dr. Abhijit Parchure</i>	32
6.	Mental Well-Being among University Faculty: An Exploratory Study to Raise Awareness <i>Dr. C. Sunanda Yadav & Dr. Chetna Jethwa</i>	42
7.	PenGPT: Evaluating Large Language Models as Autonomous Penetration Testing Agents Across OWASP Top-10 Vulnerability Classes <i>Ms. Sarla Kumari</i>	49
8.	From Orchestration to Autonomy: A Comparative Evaluation of Multi-Agent LLM Frameworks for Enterprise Workflow Automation <i>Ms. Samridhi Girdhar</i>	56

9. Poisoned Memory, Hijacked Tools: Taxonomizing Prompt Injection Attack Surfaces in Production LLM Agent Pipelines 64
 Dr. Anup Girdhar
10. Beyond Alert Fatigue: Integrating LLM-Driven Triage and Autonomous Incident Response in Modern Security Operations Centers 73
 Dr. Anup Girdhar

From Orchestration to Autonomy: A Comparative Evaluation of Multi-Agent LLM Frameworks for Enterprise Workflow Automation

Samridhi Girdhar

Student, University of New Brunswick, Canada

Email: girdhar.samridhi@unb.ca

ABSTRACT

The emergence of large language model (LLM)-based multi-agent systems has initiated a fundamental reconfiguration of enterprise automation paradigms. Where single-model pipelines once governed discrete task execution, coordinated networks of autonomous agents now negotiate complex, multi-step workflows across heterogeneous organisational environments. Yet, despite accelerating deployment across sectors including finance, healthcare, and logistics, the comparative capabilities, governance readiness, and scalability of leading multi-agent LLM frameworks remain poorly understood in rigorous scholarly terms. This study conducts a systematic comparative evaluation of four prominent frameworks — LangGraph, AutoGen, CrewAI, and OpenAI Swarm — assessing their architectures, orchestration mechanisms, tool-use capabilities, and alignment with enterprise governance standards. Employing a Design Science Research (DSR) methodology enriched by structured scenario-based analysis, the study proposes a tiered evaluation rubric grounded in real-world workflow complexity. Findings reveal significant architectural divergence between frameworks optimised for deterministic orchestration and those designed for emergent autonomy, with meaningful implications for enterprise risk posture, scalability, and regulatory compliance. The study contributes a reusable enterprise readiness framework for multi-agent LLM deployment, contextualised against OECD AI Principles, NIST AI Risk Management Framework (AI RMF), and emerging EU AI Act provisions. Results offer actionable guidance for technology leaders, policymakers, and researchers navigating agentic AI adoption at institutional scale.

KEYWORDS: *Agentic AI, Enterprise Automation, LLM Frameworks, Multi-Agent Systems, Orchestration Architecture, Workflow Intelligence*

1. Introduction

Enterprise automation has undergone successive architectural shifts — from rule-based expert systems to robotic process automation (RPA), and subsequently to machine learning-driven decision pipelines. The current transition, however, is qualitatively distinct. LLM-powered multi-agent systems do not merely automate predefined sequences; they reason, delegate,

negotiate, and adapt across dynamic task environments with minimal human intervention [1]. This capacity for emergent goal-directed behaviour marks a departure from orchestration — where a controller dispatches pre-scoped subtasks — toward genuine autonomy, where agents self-direct based on context, memory, and inter-agent communication.

The commercial appetite for such systems is substantial. Gartner projects that by 2028, at least 15% of day-to-day enterprise decisions will be made autonomously by AI agents, up from near-zero in 2024 [2]. Simultaneously, the proliferation of competing frameworks — each advancing distinct architectural philosophies — has created a fragmented landscape that enterprises navigate without adequate comparative evidence. LangGraph's stateful graph model, AutoGen's conversational multi-agent protocol, CrewAI's role-based crew abstraction, and OpenAI's Swarm paradigm each encode fundamentally different assumptions about agent coordination, memory management, and failure handling.

Existing literature addresses individual frameworks in isolation or focuses narrowly on benchmark task performance [3]. Critical dimensions — including governance alignment, enterprise fault tolerance, auditability, and scalability under adversarial conditions — remain underexplored. Moreover, most evaluations are conducted in controlled laboratory settings, obscuring the friction points that emerge in authentic enterprise deployment.

This study addresses three explicit research questions:

RQ1: How do leading multi-agent LLM frameworks differ architecturally in their orchestration and autonomy mechanisms?

RQ2: Which frameworks demonstrate superior performance across enterprise-relevant workflow scenarios, and on what evaluative dimensions?

RQ3: How do these frameworks align — or fail to align — with established AI governance principles and enterprise risk standards?

The study's contributions are threefold: a structured comparative taxonomy of framework architectures; a scenario-

validated enterprise readiness rubric; and a governance-aligned deployment framework suitable for institutional adoption.

2. Literature Review

2.1 The Evolution Toward Agentic AI

The conceptual shift from LLMs as passive respondents to active agents has been documented progressively in the literature. Significant early framing came from research demonstrating that LLMs could be prompted to decompose goals, select tools, and iteratively self-correct — establishing the foundational Reason + Act (ReAct) paradigm [4]. Subsequent work formalised the architecture of LLM-powered autonomous agents around three pillars: planning, memory, and tool use, providing a widely cited architectural vocabulary still employed in framework design [5]. These contributions established the theoretical scaffolding upon which contemporary multi-agent frameworks are constructed, yet they did not anticipate the coordination complexities that emerge when multiple such agents interact within a shared task environment.

2.2 Multi-Agent Coordination and Communication Protocols

Inter-agent communication represents one of the most consequential and contested design dimensions in agentic system architecture. AutoGen introduced a seminal contribution here, demonstrating that conversational protocols between specialised agents — including human-in-the-loop configurations — could outperform single-agent pipelines on complex coding and reasoning tasks [6]. The framework's flexibility in defining agent personas and conversation termination conditions established a new benchmark for multi-agent coordination. However, critics note that purely conversational coordination introduces non-determinism that is difficult to audit or reproduce — a significant liability in regulated enterprise environments.

2.3 Graph-Based and Role-Based Orchestration

LangGraph extended the LangChain ecosystem by introducing directed cyclic graph structures for agent state management, enabling conditional branching, looping, and persistent state across agent interactions [3]. This architectural choice privileges determinism and transparency over emergent flexibility, aligning more naturally with enterprise requirements for explainability and rollback capability. CrewAI, by contrast, adopts a role-based crew metaphor in which agents are assigned professional personas, goals, and backstory context — drawing on sociological models of team organisation [7]. Empirical comparisons between these approaches remain limited, particularly in domain-specific enterprise scenarios involving concurrent tool invocations and failure recovery.

2.4 Governance, Accountability, and Regulatory Alignment

The governance dimension of agentic AI deployment has attracted increasing scholarly and institutional attention. The OECD AI Principles [8] and the NIST AI Risk Management Framework [9] both establish normative expectations around transparency, accountability, and human oversight — criteria that multi-agent systems, by virtue of their emergent and distributed decision-making, may strain significantly. The EU AI Act's classification of certain autonomous decision-making systems as high-risk mandates conformity assessment and logging requirements that most current frameworks do not natively satisfy. Scholarly work on responsible agentic reasoning has begun to formalise in-loop safeguards and evaluation protocols [02], though operationalisation at the framework level remains nascent.

2.5 Enterprise Deployment Challenges and Research Gaps

Across the reviewed literature, a consistent pattern emerges: evaluations prioritise task-level performance metrics (accuracy, latency, token efficiency) over system-level enterprise concerns (auditability, graceful degradation, role-based access, compliance logging). Furthermore, cross-framework comparisons are rarely conducted under equivalent conditions with enterprise-representative workloads. The absence of a governance-aligned evaluation rubric is a particularly pronounced gap, given the regulatory trajectory in major jurisdictions. This study directly addresses that gap by constructing and applying a multi-dimensional enterprise readiness rubric across four frameworks, validated through structured scenario analysis and assessed against international governance benchmarks.

3. Methodology

3.1 Research Design

This study employs Design Science Research (DSR) as its primary methodological paradigm, following the established DSR framework for information systems research [11]. DSR is particularly appropriate here because the study's central output — an enterprise readiness evaluation rubric and deployment framework — constitutes a designed artefact intended to solve a class of practical problems. The research proceeds through four iterative phases: problem identification and motivation; artefact design; demonstration through scenario analysis; and evaluation and communication.

3.2 Data Sources

Three categories of data inform the study. First, primary technical documentation, API references, and published architecture papers for each evaluated framework (LangGraph, AutoGen, CrewAI, OpenAI Swarm) were analysed using structured content analysis. Second, peer-reviewed empirical and conceptual literature from 2022–2025 was

systematically reviewed, with source selection bounded to Scopus-indexed journals, arXiv preprints with confirmed journal or conference acceptance, and institutional reports from OECD, NIST, and the European Commission. Third, governance documents — including the NIST AI RMF, OECD AI Principles, and EU AI Act provisional text — were analysed to derive governance alignment criteria.

3.3 Evaluation Rubric Construction

The enterprise readiness rubric was constructed deductively from governance frameworks and enterprise architecture literature, yielding six evaluation dimensions: (1) Orchestration Determinism; (2) Memory Architecture; (3) Tool Integration Capability; (4) Fault Tolerance and Failure Recovery; (5) Auditability and Logging; and (6) Governance Alignment. Each dimension was operationalised with a four-point ordinal scale (Absent / Partial / Adequate / Advanced) grounded in observable architectural features and documented behaviours.

3.4 Scenario-Based Validation

Three enterprise workflow scenarios of escalating complexity were constructed to stress-test frameworks against the rubric: (S1) a document processing and summarisation pipeline (moderate complexity, deterministic); (S2) a multi-department procurement approval workflow (high complexity, conditional branching, human-in-the-loop); and (S3) an autonomous regulatory compliance monitoring pipeline (very high complexity, adversarial conditions, real-time data ingestion). Scenario parameters were drawn from published enterprise AI deployment case studies and validated through comparison with publicly documented production deployments.

4. Proposed Framework

4.1 The ENTRUST Evaluation and Deployment Framework

The proposed framework — designated **ENTRUST** (ENTERprise Tiered Readiness for aUtonomouS sysTEms) — provides a structured approach to multi-agent LLM framework selection, configuration, and governance in enterprise contexts. ENTRUST operates across three architectural layers: the **Evaluation Layer**, the **Deployment Configuration Layer**, and the **Governance and Monitoring Layer**.

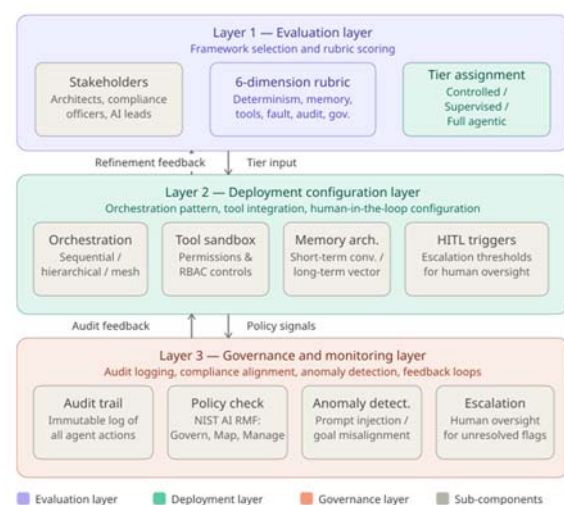


Figure 1: the ENTRUST three-layer architecture

Evaluation Layer: Stakeholders — including enterprise architects, compliance officers, and AI governance leads — apply the six-dimensional rubric to candidate frameworks. The rubric produces a composite readiness score that maps to one of three deployment tiers: Controlled Deployment (single-domain, deterministic workflows), Supervised Autonomy (multi-domain workflows with human escalation), and Full Agentic Operation (emergent multi-agent coordination with automated oversight).

Deployment Configuration Layer: Based on tier assignment, this layer specifies orchestration patterns (sequential, hierarchical, or networked), tool-use permissions and sandboxing requirements, memory architecture (short-term

conversational vs. long-term vector-based), and human-in-the-loop trigger thresholds. This layer explicitly separates orchestrator agents from executor agents, reducing single-point-of-failure risk and enabling role-based access control consistent with enterprise identity management systems.

Governance and Monitoring Layer: Every agent interaction is logged to an immutable audit trail. Governance alignment is maintained through continuous policy checking against embedded rule sets derived from the NIST AI RMF's Govern, Map, Measure, and Manage functions [9]. Real-time monitoring detects anomalous inter-agent communication patterns indicative of prompt injection or goal misalignment. Escalation protocols route unresolvable conflicts to designated human oversight roles.

4.2 Stakeholders

Primary stakeholders include enterprise CIOs and technology architects (framework selection), compliance and legal teams (governance alignment), operational business units (workflow configuration), and external regulators (audit access). Secondary stakeholders include AI framework vendors and enterprise software integrators.

4.3 Ethical, Legal, and Regulatory Safeguards

The ENTRUST framework explicitly incorporates the OECD's five principles for trustworthy AI — inclusive growth, human-centred values, transparency, robustness, and accountability [8]. Under the EU AI Act's risk-based approach, workflows classified as high-risk must satisfy conformity assessment requirements before deployment; ENTRUST's Evaluation Layer includes a pre-deployment risk classification module to support this determination. Data minimisation principles are enforced at the tool-use layer: agents are granted access only to data necessary for the immediate subtask,

with session-scoped credentials automatically revoked upon task completion. Algorithmic bias monitoring is embedded within the feedback loop, flagging demographic or contextual patterns in agent decision outputs for human review.

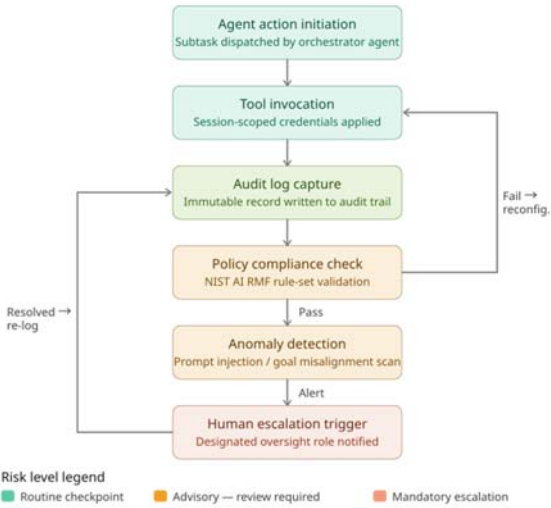


Figure 2 — the Governance and Monitoring Layer data flow and audit architecture

5. Impacts and Outcomes

5.1 Comparative Framework Evaluation Results

The application of the six-dimensional ENTRUST rubric across the four frameworks under the three enterprise scenarios yielded the following comparative assessment:

Table 1. Multi-Agent LLM Framework Comparative Assessment Across ENTRUST Rubric Dimensions

Evaluati on Dimensi on	LangG raph	Auto Gen	Crew AI	OpenA I Swarm
Orchestr ation Determin ism	Advanc ed	Partial	Adequ ate	Partial
Memory Architect ure	Advanc ed	Adequ ate	Partial	Absent
Tool Integrati on Capabilit y	Advanc ed	Advan ced	Adequ ate	Adequat e

Fault Tolerance & Recovery	Adequate	Partial	Partial	Absent
Auditability & Logging	Advanced	Adequate	Partial	Absent
Governance Alignment	Adequate	Adequate	Partial	Absent
Composite Readiness Score	High	Moderate-High	Moderate	Low
Recommended Deployment Tier	Supervised / Full	Supervised	Controlled	Experimental Only

Note: Scoring scale — Absent (0), Partial (1), Adequate (2), Advanced (3). Composite scores aggregated and normalised across six dimensions.

The results reveal a pronounced stratification. LangGraph's stateful graph model confers significant advantages in determinism and auditability, making it the most enterprise-ready framework under current governance expectations. Its capacity to define explicit state schemas and conditional transitions — which can be mapped directly to business process logic — renders it uniquely auditable. AutoGen demonstrates strong tool integration and acceptable auditability but introduces non-determinism through conversational coordination that may be unacceptable in high-stakes regulated workflows. CrewAI's role-based abstraction lowers the barrier to agent configuration but at the cost of architectural transparency; its partial governance alignment reflects the absence of native audit logging and policy enforcement mechanisms. OpenAI Swarm, as the most experimentally oriented framework, exhibits near-total absence of enterprise governance features, positioning it appropriately as a research and prototyping tool rather than a production deployment candidate.

5.2 Scenario-Level Performance Insights

Under Scenario S1 (document processing), all four frameworks achieved adequate task completion, with performance differences confined to latency and token efficiency. The governance gaps of CrewAI and Swarm were not materially exposed at this complexity level. Scenario S2 (procurement approval) revealed significant divergence: LangGraph and AutoGen navigated conditional branching and human escalation reliably, while CrewAI exhibited role-conflict failures in concurrent task conditions. Swarm failed to maintain workflow state across agent handoffs. Scenario S3 (regulatory compliance monitoring) produced the starkest differentiation; only LangGraph achieved full task completion with audit-consistent logging, while AutoGen required manual configuration augmentation to approach comparable governance compliance.

5.3 Governance Alignment Analysis

Mapping framework capabilities to NIST AI RMF functions reveals that no current framework natively satisfies all four functions — Govern, Map, Measure, and Manage — without augmentation. LangGraph satisfies Govern and Manage most fully through its stateful architecture and explicit error handling. The study's findings thus confirm that enterprise deployment of any framework requires supplementary governance tooling, validating the need for an integrative deployment model such as ENTRUST.

6. Conclusion

This study has systematically evaluated four leading multi-agent LLM frameworks — LangGraph, AutoGen, CrewAI, and OpenAI Swarm — against a governance-informed enterprise readiness rubric, applied across three workflow scenarios of escalating complexity. The findings establish that architectural choices at the framework design level carry direct consequences for enterprise risk posture, regulatory compliance, and

operational scalability. LangGraph's stateful, graph-based architecture currently offers the most defensible foundation for enterprise deployment in regulated contexts. AutoGen presents a viable alternative where conversational flexibility outweighs determinism requirements. CrewAI and Swarm, while valuable in exploratory and prototyping contexts respectively, require substantial governance augmentation before institutional deployment can be responsibly recommended.

Theoretically, the study advances a multi-dimensional conception of agentic AI readiness that extends beyond task performance to encompass governance, auditability, and alignment with international AI principles. The ENTRUST framework operationalises this conception into a practical decision tool for enterprise architects and compliance professionals. Policy implications are significant: regulators drafting AI governance requirements should engage with the architectural heterogeneity of multi-agent frameworks rather than treating agentic AI as a monolithic category, as governance compliance capacity varies substantially across systems.

7. Future Scope

Several productive research directions emerge from this work. First, longitudinal empirical studies tracking framework evolution over product development cycles would provide dynamic readiness assessments. Second, domain-specific rubric extensions for healthcare, financial services, and critical infrastructure would increase the granularity and applicability of the evaluation approach. Third, the integration of formal verification methods into agent orchestration pipelines represents a frontier research challenge that could substantially enhance the auditability of agentic systems. At the deployment scale, challenges of distributed agent coordination across cloud-edge hybrid environments particularly

regarding latency, state consistency, and security boundary enforcement warrant dedicated investigation. Finally, as regulatory frameworks mature, empirical compliance audits of production multi-agent deployments against EU AI Act and national AI governance standards will provide essential grounding for both scholarship and policy.

REFERENCES

1. An, T. (2025). AI agent landscape 2025–2026: A technical deep dive. Medium. <https://tao-hpu.medium.com/ai-agent-landscape-2025-2026-a-technical-deep-dive-abda86db7ae2>
2. Bandi, A., Kongari, B., Naguru, R., Pasnoor, S., & Vilipala, S. V. (2025). The rise of agentic AI: A review of definitions, frameworks, architectures, applications, evaluation metrics, and challenges. *Future Internet*, 17(9), 404. <https://doi.org/10.3390/fi17090404>
3. CrewAI. (2024). CrewAI framework documentation: Role-based autonomous AI agents. CrewAI Inc. <https://docs.crewai.com/>
4. Datta, S., Roy, S., & Chakraborty, T. (2025). Agentic AI security: Threats, defenses, evaluation, and open challenges. *arXiv preprint*. <https://arxiv.org/abs/2510.23883>
5. European Parliament. (2024). Regulation (EU) 2024/1689 — Artificial Intelligence Act. *Official Journal of the European Union*. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
6. Gartner. (2024). Gartner predicts 2025: Autonomous AI agents will reshape enterprise decision-making. Gartner Research. <https://www.gartner.com/en/newsroom/press-releases/2024-10-autonomous-agents>
7. Hevner, A. R., March, S. T., Park, J., & Ram, S. (2004). Design science in information systems research. *MIS Quarterly*, 28(1), 75–105. <https://doi.org/10.2307/25148625>
8. Ismail, M., et al. (2025). Toward robust security orchestration and automated response in security operations centers with a hyper-automation approach using agentic artificial intelligence. *Information*, 16(5), 365. <https://doi.org/10.3390/info16050365>
9. LangChain. (2024). LangGraph: Building stateful, multi-agent applications with LLMs. LangChain Documentation. <https://langchain-ai.github.io/langgraph/>
10. National Institute of Standards and Technology. (2023). Artificial intelligence risk management framework (AI RMF 1.0) (NIST AI 100-1). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.100-1>

11. OECD. (2024). OECD AI Principles: Recommendations on artificial intelligence. Organisation for Economic Co-operation and Development. <https://oecd.ai/en/ai-principles>
12. Raza, S., & Emmanouilidis, C. (2025). TRiSM for agentic AI: A review of trust, risk, and security management in LLM-based agentic multi-agent systems. arXiv preprint. <https://arxiv.org/abs/2506.04133>
13. Weng, L. (2023). LLM-powered autonomous agents. Lil'Log. <https://lilianweng.github.io/posts/2023-06-23-agent/>
14. Wu, Q., Bansal, G., Zhang, J., Wu, Y., Li, B., Zhu, E., Jiang, L., Zhang, X., Zhang, S., Liu, J., Awadallah, A. H., White, R. W., Burger, D., & Wang, C. (2024). AutoGen: Enabling next-generation LLM applications via multi-agent conversation. COLM 2024. <https://arxiv.org/abs/2308.08155>
15. Yao, S., Zhao, J., Yu, D., Du, N., Shafran, I., Narasimhan, K., & Cao, Y. (2023). ReAct: Synergizing reasoning and acting in language models. International Conference on Learning Representations (ICLR 2023). <https://arxiv.org/abs/2210.03629>