

REVIEWER'S REPORT

Manuscript No.: IJAR-59194

Title: Designing the AI Web: Architectural Principles for Interoperable Artificial Intelligence Systems.

Recommendation:

Accept as it is

Accept after minor revision.....

Accept after major revisionyes.....

Do not accept (*Reasons below*)

Rating	Excel.	Good	Fair	Poor
Originality			Y	
Techn. Quality			Y	
Clarity			Y	
Significance			Y	

Reviewer's ID: JPR-Dr.Shaweta Sachdeva

Detailed Reviewer's Report

1. The manuscript addresses an important and emerging topic; however, the original contribution and novelty of the proposed "AI Web" framework need to be stated much more explicitly. The authors should clearly distinguish their nine-layer architecture from existing work on distributed AI, multi-agent systems, Semantic Web, agentic AI, and Horvitz's earlier "global web of artificial intelligence" concept.
2. The manuscript claims to use a Systematic Literature Review (SLR) and mentions PRISMA guidelines, but the methodology does not provide sufficient information to qualify as a reproducible SLR. The authors should report the search dates, number of records retrieved from each database, duplicate-removal process, screening stages, final number of included studies, and reasons for exclusion.
3. A PRISMA flow diagram should be added. The current methodology states that PRISMA guidelines were followed, but no corresponding screening flow or quantitative evidence is presented.
4. The inclusion/exclusion criteria are useful but remain relatively broad. The authors should define the quality assessment criteria used to evaluate the selected studies and explain how the reliability of sources was established.
5. The proposed nine-layer AI Web architecture is the central contribution, but its design rationale needs substantially more technical justification. For each layer, the authors should explain its interfaces, inputs/outputs, dependencies, and relationship with adjacent layers.
6. The manuscript would benefit from a detailed architectural diagram showing how the nine layers interact. The current table provides a useful summary, but it does not sufficiently demonstrate the actual flow of data, context, tasks, authentication, coordination, and governance across the layers.
7. The discussion of MCP and A2A should be made more technically precise. In particular, the manuscript should distinguish tool/context interoperability from agent-to-agent communication and semantic interoperability, rather than treating these mechanisms as equivalent components of a semantic layer.
8. Several statements regarding MCP, A2A, JSON-RPC, gRPC, and semantic communication are presented as established facts but require more direct references to official protocol specifications

REVIEWER'S REPORT

and recent technical literature. The reference list should be strengthened with primary sources wherever possible.

9. The comparison of REST, GraphQL, gRPC, MCP, and A2A in Figure 2 is potentially valuable, but the qualitative scoring methodology is not explained. The basis for assigning scores such as 2, 3, 4, and 5 should be explicitly defined, and the evaluation should either be supported by measurable criteria or clearly labelled as an expert/conceptual assessment.
10. Figure 1 presents the distribution of reviewed sources by publisher, but publisher distribution does not establish literature quality or research significance. Consider replacing or supplementing this figure with more informative bibliometric information, such as publication year, research theme, protocol studied, methodology, or citation distribution.
11. The case studies are useful for illustrating the proposed architecture, but it is unclear whether they represent real-world implementations, published case studies, or hypothetical scenarios. This distinction should be explicitly stated for each case study.
12. Several case-study claims, such as “fewer diagnostic errors,” “less supply-chain delays,” and “models based on a knowledge graph increase accuracy,” require empirical evidence or appropriate references. If these are conceptual expectations rather than measured outcomes, the wording should be revised accordingly.
13. The healthcare case study requires greater technical and regulatory precision. The relationship between FHIR, privacy-preserving computation, zero-knowledge proofs, and HIPAA compliance should be explained carefully, because using a privacy-preserving technique does not by itself establish regulatory compliance.
14. The statement that knowledge-graph-based systems result in “fewer hallucinations” and improved accuracy should be supported by quantitative experimental evidence or appropriately qualified as a potential benefit rather than a demonstrated result.
15. The manuscript would be considerably stronger if the authors included quantitative evaluation metrics for interoperability. Possible metrics include context preservation, semantic fidelity, task completion rate, communication latency, token/compute overhead, failure recovery rate, security violations, and cross-platform compatibility.
16. The paper currently remains predominantly conceptual, while some portions of the abstract and discussion appear to imply empirical validation. The authors should clearly distinguish conceptual synthesis from experimentally validated findings. The abstract states that the model was “conceptually evaluated,” which should be consistently reflected throughout the manuscript.
17. The manuscript should provide at least one prototype, simulation, or proof-of-concept implementation of the proposed architecture. Even a small multi-agent implementation demonstrating MCP/A2A-based communication would substantially strengthen the contribution.
18. The proposed architecture should address failure handling and conflict resolution in greater depth. The manuscript identifies infinite loops and conflicting agent objectives as challenges, but the proposed coordination layer does not specify concrete mechanisms such as consensus, negotiation, timeout, rollback, arbitration, or escalation to humans.
19. The security discussion should be expanded beyond encryption and authentication. An interoperable AI Web should also consider prompt injection, tool poisoning, identity spoofing, privilege escalation, malicious agents, data exfiltration, supply-chain attacks, replay attacks, and compromised model/tool registries.
20. The statement that agents should retain an “auditable log of their reasoning” should be reconsidered. It would be more technically appropriate to focus on auditable decision traces, tool calls, inputs/outputs, policies, provenance, and execution logs, rather than assuming that internal model reasoning can or should be fully exposed.

International Journal of Advanced Research

Publisher's Name: Jana Publication and Research LLP

www.journalijar.com

REVIEWER'S REPORT

21. The authors should clarify the distinction between interoperability, portability, composability, compatibility, and semantic interoperability, as these terms are sometimes used interchangeably in the manuscript.
22. The research gap is relevant but could be made more rigorous by providing a comparative literature matrix showing existing frameworks, their layers, protocols, security mechanisms, semantic capabilities, governance mechanisms, and limitations. This would provide stronger evidence that the proposed nine-layer model fills a genuine gap.
23. The manuscript should discuss the performance overhead introduced by multi-agent communication. While interoperability can improve modularity, frequent agent-to-agent calls can increase latency, bandwidth consumption, token usage, and computational cost.
24. The discussion of decentralisation should be refined. The proposed AI Web is described as “fully decentralised,” yet some of the examples—including the enterprise knowledge-graph case—use centralised components. The authors should clearly define what level of decentralisation is intended and whether the architecture supports hybrid, federated, or fully decentralised deployments.
25. The manuscript should provide a stronger treatment of identity, trust, agent discovery, capability discovery, authorization, and reputation management. These are fundamental requirements for an open network in which independently operated agents interact.
26. The discussion of governance would benefit from a clearer accountability model. For example, when an autonomous workflow involving multiple agents produces an incorrect or harmful outcome, the manuscript should explain how responsibility can be assigned among the model provider, agent developer, tool provider, orchestrator, and human operator.
27. The future-direction section contains several interesting ideas, including federated intelligence, decentralised AI, AI governance, and human-AI interaction, but some statements are forward-looking and speculative. These should be clearly distinguished from conclusions supported by the reviewed literature.
28. The connection between the proposed AI Web and Artificial General Intelligence (AGI) is speculative and should either be substantially justified with supporting literature or presented more cautiously. The manuscript itself acknowledges that this remains a contested possibility.
29. The conclusion makes strong claims regarding a “genuine turning point” and the emergence of a “decentralised, global network.” These statements should be moderated unless supported by stronger empirical evidence and broader literature coverage.
30. The manuscript requires careful language editing and proofreading. There are several grammatical and stylistic issues, including awkward constructions, inconsistent terminology, missing articles, and occasional overly long sentences. Examples include “Reliable” as an architectural principle instead of “Reliability,” “routeing” versus “routing,” and inconsistent use of “AI Web,” “AI Web architecture,” and “Internet of AI Agents.”
31. The manuscript is organised into “Chapter 1,” “Chapter 2,” etc. If this is intended as a journal research article, the structure should be converted into a conventional journal format such as Introduction, Related Work/Literature Review, Methodology, Proposed Framework, Evaluation, Discussion, Limitations, and Conclusion.
32. The reference section should be checked carefully for completeness, consistency, DOI information, formatting, and correspondence between in-text citations and references. In particular, all recent protocol-related claims should have appropriate primary or authoritative references.
33. The manuscript would benefit from a dedicated “Limitations of the Proposed Framework” subsection that addresses conceptual validation, dependence on evolving standards, lack of large-scale benchmarking, closed-source model constraints, and challenges in validating interoperability

International Journal of Advanced Research

Publisher's Name: Jana Publication and Research LLP

www.journalijar.com

REVIEWER'S REPORT

across heterogeneous agents. Some limitations are currently mentioned in the introduction but need to be integrated into the discussion.

34. Overall, the manuscript has a potentially valuable conceptual contribution, but substantial strengthening is required in methodological transparency, technical precision, empirical/experimental validation, protocol analysis, security considerations, and evidence supporting the proposed architecture. The authors should consider major revision before the work can be adequately assessed for publication.