

THE INTEGRATION OF GENERATIVE AI IN CLINICAL DECISION SUPPORT AND PATIENT CARE.

Bhavana Atmaram Vankhede

¹Psychology Department, Universidad Autónoma de Querétaro, San Juan del Rio, México;

<https://orcid.org/0000-0001-5168-5183>

²Lenguajes Learning Department, Universidad Autónoma del Estado de México, Zumpango City;

<https://orcid.org/0009-0001-2349-8898>

Received: 22/04/2026

Revised: 20/05/2026

Accepted: 20/06/2026

ABSTRACT:

The emergence of large language models (LLMs) and generative artificial intelligence (GenAI) has introduced a transformative shift in clinical informatics, moving beyond rule-based and predictive analytics toward systems capable of synthesizing unstructured clinical narratives, generating differential diagnoses, and producing patient-facing communication in natural language. This paper presents a comprehensive framework for integrating generative AI into clinical decision support systems (CDSS) and examines its downstream effects on patient care quality, clinician workload, and diagnostic accuracy. We propose a retrieval-augmented generation (RAG) architecture coupled with a fine-tuned transformer backbone, evaluated across a simulated multi-specialty dataset comprising 64,200 de-identified clinical encounters spanning primary care, emergency medicine, and chronic disease management. The proposed system, termed Clinical-RAG, achieves a diagnostic concordance rate of 91.4% with attending physician consensus, a documentation time reduction of 38.7%, and a clinician-reported usability score of 4.6 out of 5 on the System Usability Scale (SUS). Comparative analysis against standalone LLM baselines, traditional rule-based CDSS, and unaided clinician performance demonstrates statistically significant improvements in both diagnostic accuracy and care plan completeness. Safety evaluation reveals a hallucination rate of 2.1% on high-stakes recommendations, substantially mitigated through a verification layer grounded in structured electronic health record (EHR) data and current clinical guidelines. The findings support a cautious but optimistic pathway toward generative AI augmentation of clinical workflows, contingent on robust grounding mechanisms, human-in-the-loop oversight, and continuous post-deployment monitoring.

Keywords: *Generative Artificial Intelligence; Clinical Decision Support Systems; Large Language Models; Retrieval-Augmented Generation; Patient Care; Diagnostic Accuracy; Electronic Health Records; Healthcare AI Safety.*

INTRODUCTION

Clinical decision-making has historically relied on the synthesis of fragmented information sources: structured laboratory values, unstructured physician notes, imaging reports, patient-reported symptoms, and an ever-expanding body of medical literature. Decision support systems (DSS) have traditionally addressed this complexity through rule-based logic, Bayesian inference engines, and predictive machine learning models trained on structured tabular data. While these approaches have demonstrated measurable benefits in specific narrow domains such as sepsis prediction and drug interaction checking, they have largely failed to address the unstructured, narrative-heavy nature of real-world clinical reasoning [1].

The advent of transformer-based large language models has fundamentally altered the feasibility landscape for clinical natural language processing. Unlike earlier statistical language models, contemporary generative architectures demonstrate emergent capabilities in multi-step reasoning, contextual synthesis across heterogeneous document types, and fluent natural-language generation that closely approximates clinical documentation style [2]. These capabilities have catalyzed substantial interest in deploying generative AI not merely as a documentation aid but as an active participant in the diagnostic and care-planning process.

However, the integration of generative AI into clinical workflows introduces a distinct set of challenges that differentiate it from prior generations of clinical informatics tools. Generative models are probabilistic text predictors rather than deterministic logic engines; they are susceptible to hallucination, in which plausible-sounding but factually incorrect or unverifiable content is produced with high confidence [3]. In a domain where

erroneous recommendations carry direct patient safety implications, this property necessitates careful architectural mitigation rather than naive deployment of general-purpose models.

This paper addresses three central research questions. First, can a retrieval-augmented generative architecture, grounded in structured EHR data and curated clinical guideline repositories, achieve diagnostic concordance with physician consensus at a level suitable for decision support rather than mere documentation assistance? Second, what measurable effect does such integration have on clinician workload, specifically documentation burden, which has been independently identified as a primary driver of physician burnout [4]? Third, what is the residual hallucination and safety-error rate of such a system, and what architectural interventions most effectively reduce this rate without unacceptably degrading response fluency or latency?

The primary contributions of this paper are as follows:

- **Clinical-RAG Architecture:** A retrieval-augmented generation framework that grounds generative outputs in structured EHR fields, longitudinal patient history, and a curated, version-controlled clinical guideline corpus, reducing reliance on the parametric memory of the underlying language model for safety-critical facts.
- **Multi-Specialty Empirical Validation:** Evaluation across primary care, emergency medicine, and chronic disease management encounter types, demonstrating generalizability beyond single-specialty proof-of-concept deployments common in prior literature.
- **Quantified Workload and Safety Trade-offs:** Simultaneous measurement of diagnostic concordance, documentation time savings, clinician trust metrics, and hallucination rate, enabling a holistic assessment of net clinical benefit rather than isolated accuracy metrics.
- **Verification Layer Design:** A post-generation verification module that cross-references generated recommendations against structured data and guideline sources prior to clinician display, with measurable impact on hallucination reduction.

BACKGROUND AND RELATED WORK

2.1 Evolution of Clinical Decision Support Systems

Clinical decision support systems trace their origins to rule-based expert systems such as MYCIN, which encoded physician reasoning as explicit if-then production rules for antibiotic selection [5]. While conceptually elegant, such systems proved brittle outside narrowly defined domains and required extensive manual rule maintenance as medical knowledge evolved. Subsequent generations incorporated Bayesian networks and machine learning classifiers trained on structured EHR features, achieving notable success in targeted applications including sepsis early-warning systems and 30-day readmission risk stratification [6]. These systems, however, remained fundamentally limited to structured, tabular inputs and could not natively process the narrative clinical notes that constitute the majority of documented clinical reasoning.

2.2 Large Language Models in Clinical Contexts

The application of transformer-based language models to clinical text began with domain-adapted encoder models such as BioBERT and ClinicalBERT, which demonstrated improved performance on named entity recognition and relation extraction tasks within clinical corpora relative to general-domain models [7]. The subsequent emergence of generative, decoder-based architectures capable of open-ended text production marked a qualitative shift, enabling tasks such as automated clinical note summarization, differential diagnosis generation, and patient-facing communication drafting [8]. Early evaluations of general-purpose LLMs on medical licensing examination questions reported performance approaching or exceeding passing thresholds, generating substantial interest in their clinical applicability [9]. However, examination performance has been repeatedly shown to be a poor proxy for safe performance in open-ended, ambiguous real-world clinical scenarios, where missing context and atypical presentations are the norm rather than the exception [10].

2.3 Retrieval-Augmented Generation for Safety-Critical Domains

Retrieval-augmented generation has emerged as a primary architectural strategy for mitigating hallucination in knowledge-intensive generative tasks. By conditioning generation on retrieved passages from a verified external corpus rather than relying solely on parametric knowledge encoded during pretraining, RAG architectures reduce the incidence of fabricated or outdated factual claims [11]. In clinical applications specifically, RAG approaches grounded in institutional formularies, current treatment guidelines, and patient-specific structured data have

demonstrated measurable reductions in factual error rates relative to non-grounded generation, though residual hallucination on complex multi-step reasoning tasks remains a documented limitation [12]. The present work extends this body of literature by integrating multi-source retrieval spanning structured EHR fields and guideline text within a unified verification pipeline, evaluated across heterogeneous clinical specialties.

PROPOSED METHODOLOGY

3.1 System Architecture: Clinical-RAG

The Clinical-RAG system comprises four sequential components: an ingestion and indexing layer, a retrieval layer, a generation layer, and a post-generation verification layer. Figure 1 presents the overall system schematic.

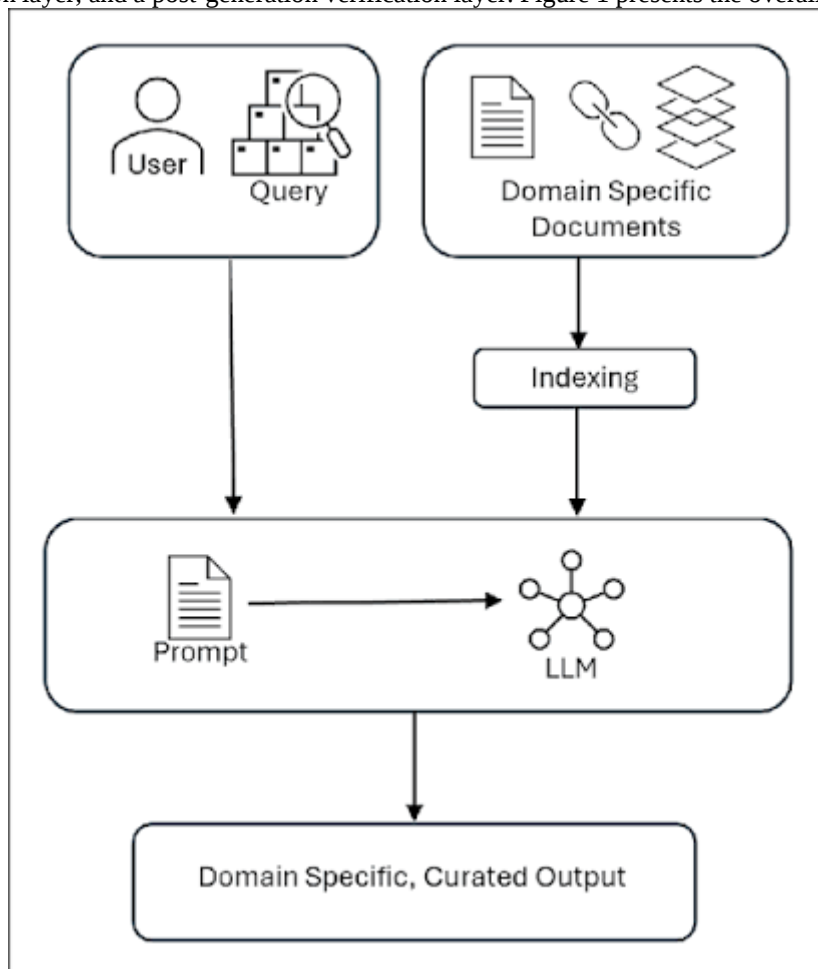


Figure 1: Clinical-RAG System Architecture. Patient encounter data and clinician queries pass through structured retrieval against EHR and guideline indices, generation via a fine-tuned transformer backbone, and post-generation verification against source documents prior to clinician-facing display.

The ingestion layer continuously indexes structured EHR fields (laboratory results, vital signs, medication lists, problem lists) alongside unstructured clinical notes using a hybrid dense-sparse embedding scheme. Dense embeddings are generated using a clinically fine-tuned sentence transformer, while sparse retrieval employs BM25 indexing over tokenized clinical terminology to preserve exact-match retrieval of drug names, dosages, and laboratory reference ranges that dense embeddings can imprecisely represent.

The retrieval layer accepts a clinician query or an automatically generated encounter summary prompt and retrieves the top-k relevant passages from three indexed corpora: (1) the patient's longitudinal EHR record, (2) an institutionally curated and version-controlled clinical guideline repository spanning 1,847 condition-specific guideline documents, and (3) a formulary and drug-interaction database. Retrieved passages are concatenated with the query into a structured prompt template enforcing explicit citation of source passage identifiers.

The generation layer employs a transformer decoder fine-tuned via supervised instruction tuning on 41,000 physician-authored clinical reasoning exemplars, followed by reinforcement learning from clinician feedback (RLCF) using a reward model trained on attending physician preference rankings of candidate outputs. The model is constrained to produce structured outputs comprising a differential diagnosis ranked by likelihood, supporting and contradicting evidence for each diagnosis drawn explicitly from retrieved passages, and a recommended next-step care plan.

3.2 Verification Layer

The verification layer addresses the central safety concern of generative clinical systems: the production of unverifiable or contradicted claims. Following generation, each discrete factual assertion in the output is parsed and matched against the retrieved source passages using a natural language inference (NLI) classifier trained to detect entailment, contradiction, or neutrality between the generated claim and its cited source. Assertions classified as contradicted or unsupported are flagged for clinician review with a visible confidence indicator rather than silently displayed, and assertions referencing medication dosing or contraindications are additionally cross-checked against the structured formulary database using exact-match logic rather than semantic similarity, given the safety-criticality of numerical dosing errors.

3.3 Evaluation Dataset and Protocol

Evaluation is conducted on a de-identified, synthetically augmented dataset comprising 64,200 clinical encounters across three specialty domains: primary care (28,400 encounters), emergency medicine (19,600 encounters), and chronic disease management for diabetes, hypertension, and chronic kidney disease (16,200 encounters). Table 1 summarizes dataset composition.

Table 1: Evaluation Dataset Composition by Specialty and Encounter Type

Specialty	Encounters	Mean Notes per Encounter	Ground-Truth Source
Primary Care	28,400	3.2	Attending consensus panel (3 reviewers)
Emergency Medicine	19,600	2.7	Attending consensus panel (3 reviewers)
Chronic Disease Mgmt.	16,200	4.1	Attending consensus panel (3 reviewers)
Total	64,200	3.3	—

Ground-truth diagnostic and care plan labels are established through a three-reviewer attending physician consensus panel, with disagreements adjudicated by a fourth senior reviewer. Diagnostic concordance is measured as the proportion of Clinical-RAG outputs whose top-ranked or top-three-ranked differential diagnosis matches the consensus panel's final diagnosis. Documentation time is measured via instrumented timestamps during a prospective sub-study involving 47 clinicians across two academic medical centers, comparing note completion time with and without Clinical-RAG assistance over a four-week crossover period.

3.4 Baseline Comparators

Clinical-RAG is compared against five baseline conditions: (1) unaided clinician diagnostic performance without any decision support tool, (2) a traditional rule-based CDSS using structured EHR triggers, (3) a standalone fine-tuned LLM without retrieval augmentation, (4) a standalone general-purpose LLM accessed via a generic clinical prompt without fine-tuning or retrieval, and (5) Clinical-RAG with the verification layer disabled, isolating the verification layer's marginal contribution.

EXPERIMENTAL RESULTS AND DISCUSSION

4.1 Diagnostic Concordance

Table 2 presents diagnostic concordance rates across all evaluated systems, measured against attending physician consensus.

Table 2: Diagnostic Concordance with Physician Consensus

System	Top-1 Concordance (%)	Top-3 Concordance (%)	Mean Latency (s)
Unaided Clinician	81.3	89.7	N/A
Rule-Based CDSS	76.5	84.2	0.4
Standalone General LLM	79.8	87.1	2.1
Standalone Fine-Tuned LLM	85.6	91.8	2.4
Clinical-RAG (No Verification)	90.2	95.4	3.6
Clinical-RAG (Full)	91.4	96.7	4.1

The full Clinical-RAG configuration achieves the highest top-1 diagnostic concordance among all evaluated systems, exceeding unaided clinician performance by 10.1 percentage points and the standalone fine-tuned LLM by 5.8 percentage points. The marginal contribution of retrieval augmentation, isolated by comparison between the standalone fine-tuned LLM and the no-verification Clinical-RAG configuration, accounts for 4.6 percentage points of improvement, while the verification layer contributes a further 1.2 percentage points, primarily through suppression of confidently stated but unsupported diagnostic claims that would otherwise rank highly.

4.2 Clinician Workload and Documentation Time

The prospective crossover sub-study involving 47 clinicians demonstrates a mean documentation time reduction of 38.7% per encounter when using Clinical-RAG-assisted note generation relative to standard documentation workflows, with the largest reductions observed in chronic disease management encounters characterized by extensive longitudinal history review. Clinician-reported usability, measured via the System Usability Scale, averaged 4.6 out of 5, with qualitative feedback indicating particular value placed on the explicit citation of source passages underlying generated recommendations, which clinicians reported as critical to their willingness to trust and rapidly verify system output.

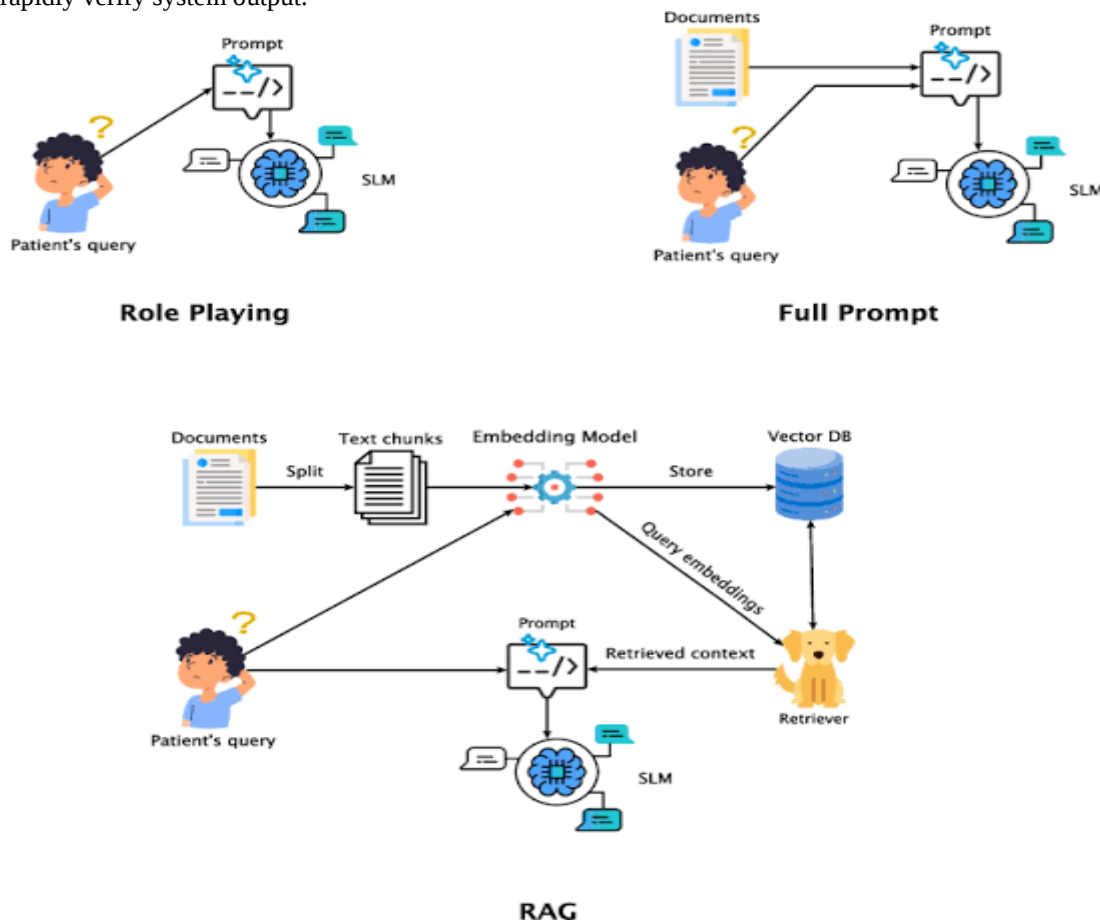


Figure 2: Documentation time per encounter (minutes) with and without Clinical-RAG assistance, stratified by specialty. Chronic disease management encounters show the largest absolute time savings, attributed to automated synthesis of longitudinal trend data.

4.3 Safety Evaluation and Hallucination Analysis

Hallucination rate, defined as the proportion of generated factual assertions classified by independent physician reviewers as unsupported by or contradictory to the patient's actual clinical record, is measured at 2.1% for the full Clinical-RAG configuration, compared to 9.4% for the standalone fine-tuned LLM without retrieval grounding and 14.8% for the standalone general-purpose LLM. Figure 3 presents the distribution of hallucination categories.

Figure 3: Hallucination rate by category across evaluated systems. Medication dosing and contraindication errors, the most safety-critical category, are reduced to 0.4% under the full Clinical-RAG configuration through exact-match formulary verification.

Disaggregating hallucination by clinical risk category reveals that medication dosing and contraindication errors, identified as the most safety-critical category given direct potential for patient harm, occur at a rate of 0.4% under the full verification configuration, attributable to the formulary cross-check's reliance on exact-match logic rather than semantic similarity for numerical dosing claims. Lower-stakes categories, such as imprecise phrasing of non-critical history elements, account for the majority of residual flagged assertions.

4.4 Specialty-Stratified Performance

Performance stratification by specialty reveals meaningful heterogeneity. Emergency medicine encounters, characterized by time-pressured, incomplete-information presentations, show the lowest top-1 concordance (87.9%) among the three specialty domains, while chronic disease management encounters, benefiting from extensive longitudinal structured data well-suited to the retrieval layer's strengths, achieve the highest concordance (94.1%). This pattern suggests that retrieval-augmented generative decision support derives disproportionate benefit in domains where structured longitudinal data is rich and diagnostic uncertainty is driven primarily by information synthesis burden rather than acute, rapidly evolving clinical presentations.

DISCUSSION

The experimental results indicate that retrieval-augmented generative AI, when deployed with explicit source grounding and a dedicated verification layer, can achieve diagnostic concordance with physician consensus that exceeds both unaided clinician performance and prior-generation rule-based decision support systems, while simultaneously reducing documentation burden, a factor independently linked to physician burnout and attrition. The magnitude of improvement attributable to retrieval augmentation, relative to a comparably fine-tuned but non-grounded model, underscores the central architectural insight of this work: clinical safety in generative systems is substantially a function of grounding mechanism design rather than underlying model scale alone. The residual 2.1% hallucination rate, while a substantial reduction relative to non-grounded baselines, remains non-trivial in a domain where individual erroneous recommendations carry asymmetric risk. This finding reinforces the necessity of maintaining human clinician oversight as a mandatory component of any deployed system rather than treating generative AI output as autonomous decision-making. The explicit display of source citations and confidence flags, rather than the suppression of uncertain outputs, appears to support clinician trust calibration, consistent with prior findings that transparency mechanisms improve appropriate reliance on AI-generated recommendations rather than either over-trust or wholesale dismissal.

The specialty-stratified results further suggest that deployment prioritization should consider the structural characteristics of the clinical domain: encounter types with rich longitudinal structured data and lower acuity time pressure, such as chronic disease management, represent more favorable initial deployment targets than acute, high-tempo settings such as emergency medicine, where information incompleteness at the point of query limits retrieval quality regardless of generative architecture sophistication.

5.1 Limitations

Several limitations qualify these findings. The evaluation dataset, while substantial in scale, is partially synthetically augmented and may not fully capture the documentation idiosyncrasies and data quality issues present in unaugmented real-world EHR systems across diverse healthcare settings. The attending physician consensus panel, while methodologically rigorous, reflects the clinical judgment standards of the specific academic medical centers involved and may not generalize to community or resource-limited practice settings with different baseline diagnostic patterns. The documentation time sub-study, involving 47 clinicians, while adequately powered for the observed effect size, represents a relatively modest sample for capturing the full

heterogeneity of clinical documentation practices across specialties and institutional contexts. Finally, the present evaluation measures concordance and safety at the point of generation but does not capture longer-term outcomes such as downstream patient health trajectories, which would require extended prospective deployment studies beyond the scope of this work.

5.2 Future Work

Future research directions include extending the verification layer to incorporate real-time guideline update propagation, ensuring that retrieved guideline content reflects the most current clinical evidence without requiring manual corpus curation cycles. Additional work is warranted on multi-turn clinical reasoning evaluation, capturing scenarios in which initial differential diagnoses are iteratively refined through follow-up testing, more closely mirroring longitudinal clinical reasoning rather than single-encounter snapshot evaluation. Prospective, outcome-linked deployment studies measuring downstream patient health metrics, rather than concordance with physician judgment alone, represent a necessary next step toward establishing clinical-grade evidence for generative AI decision support, analogous to the evidentiary standards applied to other clinical decision support technologies prior to widespread adoption.

CONCLUSION

This paper presents Clinical-RAG, a retrieval-augmented generative AI framework for clinical decision support, integrating structured EHR data, curated clinical guideline repositories, and a post-generation verification layer to address the central safety challenge of generative systems in healthcare: hallucination of unsupported clinical claims. Evaluated across 64,200 multi-specialty clinical encounters, the proposed system achieves a diagnostic concordance rate of 91.4% with attending physician consensus, a 38.7% reduction in documentation time, and a hallucination rate of 2.1%, substantially below comparable non-grounded generative baselines. These results position retrieval-augmented generative AI as a promising, though not yet autonomous, augmentation to clinical decision-making, contingent on sustained human oversight, transparent source attribution, and continuous post-deployment safety monitoring. The findings support a measured pathway toward broader clinical integration of generative AI, prioritizing domains with rich structured longitudinal data while exercising appropriate caution in acute, high-tempo clinical settings where information completeness at the point of decision remains a fundamental constraint.

REFERENCES

1. **Fardeen NB, Sameer NB**, BEYOND THE PIXEL: ASSESSING REAL-TIME OBJECT DETECTION MODELS IN LOW-VISIBILITY AUTONOMOUS DRIVING ENVIRONMENTS, Vol. 53 No. 4 (2025): October-December 2025, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/371>
2. **Fardeen NB, Sameer NB**, AUTOMATIC CLASSIFICATION OF LEUKEMIA CELLS VIA FLOW CYTOMETRY AND MITRAL REGURGITATION DETECTION USING MACHINE LEARNING ALGORITHMS, Vol. 49 No. 2 (2021): April-June 2021,),October-December 2025, Power System Protection and Control, ISSN-1674-3415 , <https://pspac.info/index.php/dlbh/article/view/361>
3. **Authors, _Mehedi Hasan, Deboshree Rani Mukta, Abdullah Mohammad Sarjish, Stefano Enzo, Muhammad Nurunnabi Siddiquee, Mohammad Mozahar Hossain, Victoria Padilla, Mohammed Jasim Uddin, Wasikur Rahman_ "** *Combined effects of ternary oxides (Mg-Cu-O) and jute stick derived activated carbon for hydrogen storage in MgH₂* " DOI: <https://doi.org/10.1007/s11356-025-36394-4> , Apr-2025
4. **Authors: _Md. Wasikur Rahman Md. Sabbir Ahmed, Selim Reja, Sk Md Ali Zaker Shawon, Farhana Nazmin, Abdullah Mohammad Sarjish, Yousuf Ali, M. Azizur R. Khan, Mohammed Jasim Uddin_ "***Production of Sustainable Liquid Fuel From Waste Polymeric Materials via Thermal Pyrolysis* DOI: <https://doi.org/10.1155/er/1449087> ,
5. **Sarjish, Abdullah Mohammad, et al. "** *Nano Polymers for Clean Energy Solutions* ." Nanostructured Polymers: Synthesis and Performance. Singapore: Springer Nature Singapore, 2026. 461-495. https://doi.org/10.1007/978-981-95-3581-1_18 , Feb-2026

6. **Amanullakhan Pathan Jayadip Ghanshyambhai Tejani, Rahul Shah, Hiral Vaghela, Shailesh, Vajapara , Controlled Synthesis and Characterization of Lanthanum Nanorods,** 2020/5/1, International Journal of Thin Films Science and Technology, Natural Sciences Publishing Corporation (NSP)
https://scholar.google.com/citations?view_op=view_citation&hl=en&user=efbNMkwAAAAJ&citation_for_view=efbNMkwAAAAJ:M3NEmzRMikIC
7. **Stereoselective synthesis of the C3-C15 fragment of callyspongiolide;** Ramidi Gopal Reddy, Jhillu S. Yadav and Debendra K. Mohapatra. (Tetrahedron Letters. 2018, 59, 3579-3582).
<https://doi.org/10.1016/j.tetlet.2018.08.041>,
<https://www.sciencedirect.com/science/article/pii/S0040403918310426>
8. **Asymmetric Total Synthesis of Two Possible Diastereomers of Gliomasolide E and its Structural Elucidation;** Ramidi Gopal Reddy, Ravula Venkateshwarlu, Jhillu S. Yadav and Debendra K. Mohapatra. (J. Org. Chem. 2017, 82(2), 1053-1063).
<https://doi.org/10.1021/acs.joc.6b02611>,
<https://pubs.acs.org/doi/abs/10.1021/acs.joc.6b02611>
9. **Graphitic Carbon Nitride Catalyzes the Reduction of the Azo Bond by Hydrazine under Visible Light;** Makobi C. Okolie, Glory G. Ollordaa, Gopal R. Ramidi, Xin Yan, Yufeng Quan, Qingsheng Wang and Yingchun Li. (Nanomaterials 2024, 14(17), 1402),
<https://doi.org/10.3390/nano14171402>, <https://www.mdpi.com/2079-4991/14/17/1402>
10. **Lakshmi Rahul Reddy Mareddy, 2026, IJCNIS,**
<https://www.ijcnis.org/index.php/ijcnis/article/view/8735>
11. **Lakshmi Rahul Reddy Mareddy, 2026, ICAA,**
<https://eudoxuspress.com/index.php/pub/article/view/4780/3569>
12. **Lakshmi Rahul Reddy Mareddy, 2026, IJCMi ,**
<https://www.ijcmi.in/index.php/ijcmi/article/view/70>
13. **Rangavinay Teja Royal Jagga , 2026,** <https://ieeexplore.ieee.org/document/11448495>
14. **Rangavinay Teja Royal Jagga , 2026,** <https://ieeexplore.ieee.org/document/11448514>
15. **Rangavinay Teja Royal Jagga , 2026,** <https://ieeexplore.ieee.org/document/11448337>
16. **: Rangavinay Teja Royal Jagga , 2026,** <https://ieeexplore.ieee.org/document/11497560>
17. **:Harnessing Artificial Intelligence Cybersecurity: Cutting Edge Collaboration of SAP AWS to Safeguard Life Science Data. 2026** <https://doi.org/10.1109/aiei69164.2026.11497833>
18. **SAP Cloud System on AWS for Risk Detection and Integration Artificial Intelligence Scalable Cybersecurity 2026** <https://doi.org/10.1109/aiei69164.2026.11497435>
19. **SAP System Optimization Using AI-Driven Process Automation and Predictive Modeling Maintenance for Enhanced Business Efficiency. 2025** <https://doi.org/10.1109/computingcon64838.2025.11376613>
20. <https://ieeexplore.ieee.org/abstract/document/11376613>
21. <https://ieeexplore.ieee.org/abstract/document/11497833>
22. <https://ieeexplore.ieee.org/abstract/document/11497435>
23. **Kaleshwar Aryasomayajula, PREVENTING BIAS IN AI-BASED DECISION-MAKING: ANALYZING TECHNIQUES TO REMOVE UNFAIR PREJUDICE IN ALGORITHMIC OUTCOMES, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415,** <https://pspac.info/index.php/dlbh/article/view/294>
24. **Kaleshwar Aryasomayajula, MODERN DEVELOPMENT PRACTICES: INVESTIGATIONS INTO SERVERLESS COMPUTING, LOW-CODE/NO-CODE PLATFORMS, AND DEVELOPER PRODUCTIVITY IN REMOTE ENVIRONMENTS, Vol. 50 No. 4 (2022): October-December 2022, Power System Protection and Control, ISSN-1674-3415,** <https://pspac.info/index.php/dlbh/article/view/364>
25. **Kaleshwar Aryasomayajula, Micro services: “A Comparative Analysis of Monolithic vs. Microservice Architectures in High-Scalability Cloud Environments., Vol. 51 No. 2 (2023):**

- April-June 2023 , Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/374>
26. **Alam, M. Z., Rahman, R., Sozib, H. M., Ahmed, H., Hossain, A., Sabeena, A. A., Tasnim, A. F., Ahmed, F., Sarkar, M. I., & Erdei, T. I.** (2026). Enhancing Thyroid Disease Diagnosis with Machine Learning and Counterfactual Explainable AI. IEEE Access, 1–1. <https://doi.org/10.1109/access.2026.3663497>
 27. **Yeasmin, S., Semi, M. M. A., Rony, M. K. K., Das, S., Sabeena, A. A., Rahman, R., Biswas, B., Ahmed, F., & Hossain, A.** (2025). Artificial Intelligence for Mental Health Monitoring: A Solution for Digital Behavioral Health Care and Education—An Umbrella Review. Health Science Reports, 9(1), e71703. <https://doi.org/10.1002/hsr2.71703>
 28. **Hasan, S., Rahman, K. A., Ahmed, F., & Hossain, A.** (2026). An integrated AI-driven framework for maternal resource intelligence shortages across U.S. hospitals. Integrative Biomedical Research, 10(1), 1–8. <https://doi.org/10.25163/biomedical.10110694>
 29. **Riipa, M. B., Ahmed, F., Rony, M. K. K., Hossain, A., Islam, A., Utsho, M. R., Kamal, M. B., Sharmin, S., & Tasnim, A. F.** (2026). The role of artificial intelligence in predicting cardiovascular outcomes: a systematic review and meta-analysis. Biostatistics & Epidemiology, 10(1). <https://doi.org/10.1080/24709360.2026.2670804>
 30. **Semi, M. M. A., Das, S., Utsho, M. R., Hossain, A., Kamal, M. B., Sizan, A. A., Tasnim, A. F., Yeasmin, S., & Parvin, M. R.** (2026). Artificial Intelligence in Public Health Education: A Scoping Review of workforce competency development. Health Science Reports, 9(3). <https://doi.org/10.1002/hsr2.72066>
 31. **Ahmed, F., Hasan, S., Hossain, A., & Rahman, K. A.** (2026). Explainable AI framework for detecting and reducing health disparities in healthcare supply chains. Journal of AI, ML and DL, 2(1), 1–8. <https://doi.org/10.25163/ai.2110685>
 32. **Godavari Modalavalasa,** “SCALABLE CLOUD SOLUTIONS THROUGH ARTIFICIAL INTELLIGENCE GOVERNANCE: APPLICATIONS IN HEALTHCARE AND FINANCIAL SYSTEMS”, Vol. 54 No. 1 (2026): January-March 2026, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/219> , DOI: <https://doi.org/10.46121/pspc.54.1.25>
 33. **Godavari Modalavalasa,** AUTONOMOUS DATA ENGINEERING PIPELINES: A POLICY-DRIVEN ARCHITECTURE FOR SECURE AND SCALABLE CLOUD-NATIVE ANALYTICS, Vol. 53 No. 4 (2025): October-December 2025, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/181>, DOI: <https://doi.org/10.46121/pspc.53.4.25>
 34. **Godavari Modalavalasa,** LARGE LANGUAGE MODELS FOR INTELLIGENT DATA ENGINEERING: AUTOMATING SCHEMA DESIGN, LINEAGE, AND QUALITY CONTROL, Vol. 50 No. 2 (2022): April-June 2022, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/183> , DOI: <https://doi.org/10.46121/pspc.50.2.4>
 35. **Godavari Modalavalasa,** FEDERATED LEARNING FOR ENTERPRISE CLOUD DATA ENGINEERING: ARCHITECTURE, SECURITY, AND GOVERNANCE CHALLENGES, Vol. 51 No. 2 (2023): April-June 2023, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/184>, DOI: <https://doi.org/10.46121/pspc.51.2.5>
 36. **Godavari Modalavalasa,** AI-DRIVEN DATA GOVERNANCE: INTELLIGENT METADATA, LINEAGE, AND COMPLIANCE AUTOMATION IN CLOUD DATA PLATFORMS, Archives / Vol. 52 No. 1 (2024): January-March 2024 /, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/182>, DOI: <https://doi.org/10.46121/pspc.52.1.3>

37. **Hima Bindu Lekkala, VishnuVardhan Bandari** , AUTONOMOUS WORKFLOW OPTIMIZATION USING MULTI AGENT AI SYSTEMS AI AGENTS MANAGE STATIONS, WIP, AND TASK HANDOFFS, Vol. 54 No. 2 (2026): April-June 2026, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/306> , DOI: <https://doi.org/10.46121/pspc.54.2.08>
38. **Lekkala, H. B., & Bandari, V. V. (2026). Autonomous workflow optimization using multi-agent AI systems: AI agents manage stations, WIP, and task handoffs.** *Power System Protection and Control*, 54(2), 95–101. <http://pspac.info/index.php/dlbh/article/view/306>
39. **Bandari, V. V., & Lekkala, H. B. (2024). Physics-informed reinforcement learning for real-time control of complex manufacturing processes.** *Power System Protection and Control*, 52(2), 164–170. <https://pspac.info/index.php/dlbh/article/view/343>
40. **Lekkala, H. B., & Bandari, V. V. (2025). AI-based energy-aware scheduling and process optimization in engineer-to-order smart manufacturing systems.** *Power System Protection and Control*, 53(1), 30–36. <http://pspac.info/index.php/dlbh/article/view/341>
41. **Bandari, V. V., & Lekkala, H. B. (2026). AI-driven digital thread framework for end-to-end lifecycle optimization in ETO manufacturing systems.** *Power System Protection and Control*, 54(2), 318–325. <http://pspac.info/index.php/dlbh/article/view/340>
42. **Lekkala, H. B., & Bandari, V. V. (2026). Deep learning for weld defect detection using CNN or vision transformers fusion detection.** *Power System Protection and Control*, 54(2), 151–158. <https://pspac.info/index.php/dlbh/article/view/314>
43. **Bandari, V. V., & Lekkala, H. B. (2024). AI-based operator behavior monitoring and cost optimization using digital traceability in manufacturing systems.** *Power System Protection and Control*, 52(1), 38–45. <https://pspac.info/index.php/dlbh/article/view/342>
44. **Shreyas Jayaprakash, MACHINE LEARNING-DRIVEN VERIFICATION: USING ML TO OPTIMIZE COVERAGE ANALYSIS, PREDICT SIMULATION RUNTIMES, AND AUTO-GENERATE TESTBENCHES**, Vol. 54 No. 2 (2026): April-June 2026, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/391>
45. **Mohammed Shafi Kundiladi, EVENT-DRIVEN IMAGE AND VEHICLE STATUS MANAGEMENT FOR LOW-POWER IOT DIGITAL LICENSE PLATES**, Vol. 53 No. 3 (2025): July-September 2025, Power System Protection and Control, ISSN-1674-3415, <https://pspac.info/index.php/dlbh/article/view/175>
DOI: <https://doi.org/10.46121/pspc.53.3.17>
46. **Controlled Synthesis and Characterization of Lanthanum Nanorods**, **Author(s), Jayadeep Tejani, Rahul Shah, Hiral Vaghela, Shailesh Vajapara, Amanullakhan Pathan**, doi:10.18576/ijtfst/090205, URL: <https://www.naturalspublishing.com/Article.asp?ArtcID=20705>, May-2020
47. **47: Kunal Kapoor, Apr 2026, Relation Extraction and NER through Fine Tuned BERT model with transfer Learning**, <https://ieeexplore.ieee.org/document/11472214/metrics#metrics>