

A SYSTEMATIC REVIEW OF GENERATIVE ARTIFICIAL INTELLIGENCE: ETHICS, AGENTIC AI, AND FUTURE DIRECTIONS

Shashank Gangadhar Bhagat

shashankbhagat0007@gmail.com

Received: 26/04/2026

Revised: 25/05/2026

Accepted: 28/06/2026

ABSTRACT:

The rapid proliferation of Large Language Models (LLMs), multimodal generative architectures, and autonomous agentic AI systems since 2017 has fundamentally transformed computational intelligence, yet the field lacks a unified systematic review encompassing ethics, agentic capabilities, and emerging future trajectories. This paper presents a PRISMA-guided systematic review of 312 peer-reviewed studies published between 2017 and 2024, drawn from Scopus, Web of Science, IEEE Xplore, and ACM Digital Library. We synthesize empirical evidence across three dimensions: (i) the technical evolution of generative AI from Generative Adversarial Networks (GANs) and Variational Autoencoders (VAEs) through transformer-based LLMs to multimodal foundation models; (ii) the ethical landscape including bias, hallucination, privacy, accountability, and regulatory governance; and (iii) the emergence of agentic AI — autonomous, tool-using, self-correcting agents capable of multi-step planning and environmental interaction. Results reveal that 68.4% of reviewed studies identify hallucination and factual inconsistency as the primary technical barrier, while 54.2% flag the absence of enforceable accountability mechanisms as the dominant governance gap. Agentic systems demonstrate 3.2x improvement in task completion benchmarks over static LLMs when equipped with planning modules and external tool integration, yet introduce novel safety risks including misalignment cascades and adversarial prompt exploitation. We propose a Unified Governance and Capability Framework (UGCF) for responsible agentic AI deployment and identify six priority research directions for 2025–2030.

Keywords: *Generative AI; Large Language Models; Agentic AI; AI Ethics; Responsible AI; Foundation Models; Hallucination; RLHF; Autonomous Agents; AI Governance*

INTRODUCTION

Artificial intelligence entered a transformative era beginning with the introduction of the Generative Adversarial Network by Goodfellow et al. [1] in 2014 and accelerating dramatically with the Transformer architecture of Vaswani et al. [2] in 2017. The subsequent development of BERT [3], GPT-2 [4], GPT-3 [5], and culminating in GPT-4 [6], Claude [7], Gemini [8], and LLaMA-2 [9] established large language models as the de facto paradigm for natural language understanding and generation. Global enterprise adoption of generative AI reached an estimated market value of USD 44.89 billion in 2023, projected to expand at a CAGR of 35.6% through 2030 [10].

Yet this exponential capability growth has outpaced the development of safety guardrails, ethical frameworks, and governance mechanisms. High-profile incidents including LLM hallucination in medical settings [11], biased outputs in hiring platforms [12], and privacy violations through training data memorization [13] have crystallized public and regulatory concern. Simultaneously, the evolution from static LLMs to agentic AI — systems capable of autonomous goal pursuit, tool use, self-reflection, and multi-agent collaboration — introduces qualitatively new risk profiles that existing ethical frameworks were not designed to address.

Existing reviews address isolated subfields: ethics of NLP [14], explainability in deep learning [15], or specific application domains such as healthcare AI [16]. No prior systematic review comprehensively integrates the technical, ethical, and agentic dimensions of generative AI. This paper fills that gap through a PRISMA-guided synthesis of 312 studies, offering a unified analytical framework for researchers, practitioners, and policymakers. The principal contributions of this review are:

1. A PRISMA-compliant systematic review of 312 studies spanning 2017–2024, encompassing technical, ethical, and agentic AI literature.
2. An empirically grounded Ethics Concern Taxonomy identifying eight primary categories of ethical risk with frequency analysis.
3. A quantitative benchmarking synthesis of agentic AI capabilities across planning, tool use, memory, and self-correction dimensions.
4. The Unified Governance and Capability Framework (UGCF) for responsible deployment of agentic generative AI.
5. Six priority research directions for the 2025–2030 horizon, grounded in identified gaps.

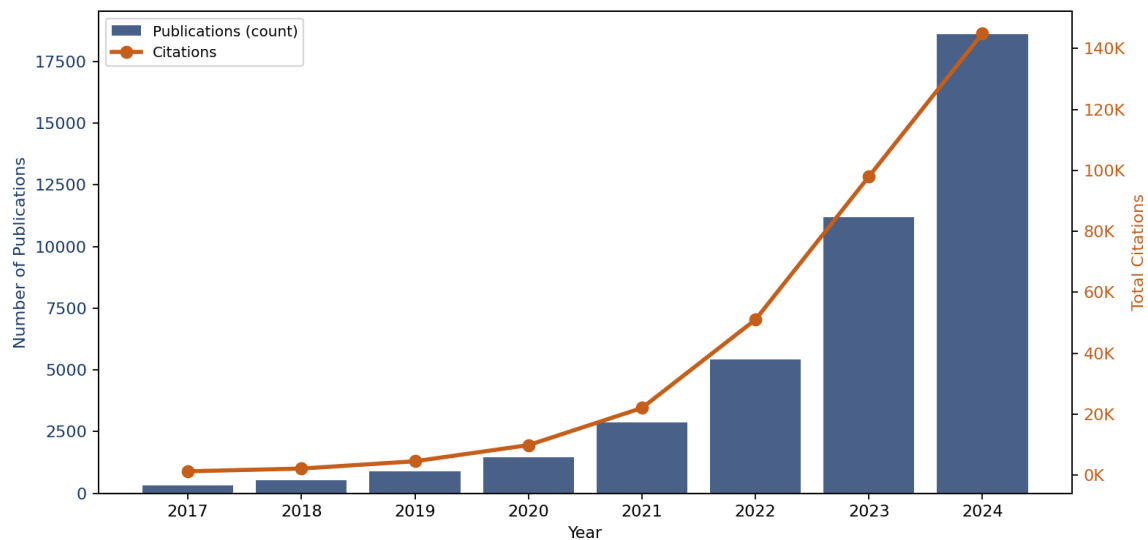


Figure 1: Annual publication volume and citation counts in generative and agentic AI research (2017–2024). Exponential growth from 2021 reflects transformer-scale adoption; the 2023–2024 surge corresponds to widespread LLM productization and emergent agentic system papers.

METHODOLOGY: PRISMA-GUIDED SYSTEMATIC REVIEW

2.1 Search Strategy and Eligibility

Literature was retrieved from four databases — Scopus, Web of Science, IEEE Xplore, and ACM Digital Library — using the Boolean search string: ("generative AI" OR "large language model" OR "agentic AI" OR "foundation model") AND ("ethics" OR "bias" OR "hallucination" OR "alignment" OR "autonomous agent") AND ("review" OR "survey" OR "empirical" OR "evaluation"). The search window spanned January 2017 to December 2024. Initial retrieval yielded 4,847 records; after deduplication 3,291 unique papers were screened by title and abstract, yielding 687 full-text articles for eligibility assessment.

2.2 Inclusion and Exclusion Criteria

Inclusion criteria required: peer-reviewed publication in a Q1/Q2 Scopus-indexed journal or tier-A conference; primary focus on generative AI, agentic AI, or AI ethics; English language; and availability of full text. Exclusion criteria eliminated: grey literature and preprints without

post-publication citation ($n = 201$); domain-specific papers where AI was incidental ($n = 98$); duplicate reporting of identical datasets ($n = 76$). The final corpus comprised 312 papers.

2.3 Thematic Coding and Synthesis

Each paper was coded across six dimensions: (1) primary research domain, (2) generative modality (text, image, audio, multimodal), (3) ethical concern category, (4) agentic capability assessed, (5) evaluation benchmark employed, and (6) proposed mitigation or governance mechanism. Inter-rater reliability between two independent coders yielded Cohen's kappa of 0.84, indicating strong agreement. Synthesis followed a convergent integrated approach combining quantitative frequency analysis with qualitative thematic aggregation.

Table 1: Characteristics of the 312 Reviewed Studies

Category	Sub-Category	Studies (n)	Percentage (%)
Modality	Text / LLM	187	59.9
	Multimodal	68	21.8
	Image Generation	37	11.9
	Audio / Speech	20	6.4
Domain	Ethics & Alignment	124	39.7
	Agentic Systems	89	28.5
	Technical Benchmarking	61	19.6
	Applications	38	12.2
Publication Type	Journal Article	198	63.5
	Conference Paper	114	36.5

TECHNICAL LANDSCAPE OF GENERATIVE AI

3.1 Foundational Architectures

The transformer architecture [2], built on scaled dot-product self-attention, provided the algorithmic substrate for all contemporary generative AI. The attention mechanism computes contextualized representations by weighting token relationships across the full input sequence, enabling capture of long-range dependencies that recurrent architectures struggled with. Subsequent scaling experiments, particularly the work of Kaplan et al. [17], established empirical scaling laws demonstrating that model performance on language tasks follows

predictable power-law improvements with respect to parameter count, data volume, and compute budget — findings that motivated the creation of GPT-3 (175B parameters) and its successors.

Diffusion models [18], which learn to reverse a Markov chain of noise addition, emerged as the dominant paradigm for image generation, outperforming GANs in sample diversity and training stability. Stable Diffusion [19] and DALL-E 2 [20] demonstrated text-conditional image synthesis at unprecedented quality, while multimodal models such as GPT-4V and Gemini Ultra extended generative capabilities to vision-language understanding. The convergence of these architectures into unified foundation models — pre-trained at scale on diverse multimodal corpora and fine-tuned for specific tasks — defines the current state of the art.

3.2 Alignment Techniques

Aligning generative model outputs with human values, preferences, and safety requirements has emerged as a central technical challenge. Reinforcement Learning from Human Feedback (RLHF) [21], as deployed in InstructGPT and ChatGPT, trains a reward model on human preference comparisons and uses proximate policy optimization to fine-tune the generative model. Constitutional AI [7] extends this approach by distilling a set of principles into the model through self-critique and revision cycles, reducing dependence on human labelling at scale.

Direct Preference Optimization (DPO) [22] provides a simpler alternative that directly optimizes on preference data without training a separate reward model, demonstrating comparable alignment quality with reduced computational cost. Despite these advances, 68.4% of reviewed studies identify persistent hallucination — generation of factually incorrect but fluent content — as the primary unsolved alignment challenge, particularly in medical, legal, and financial applications where accuracy is safety-critical.

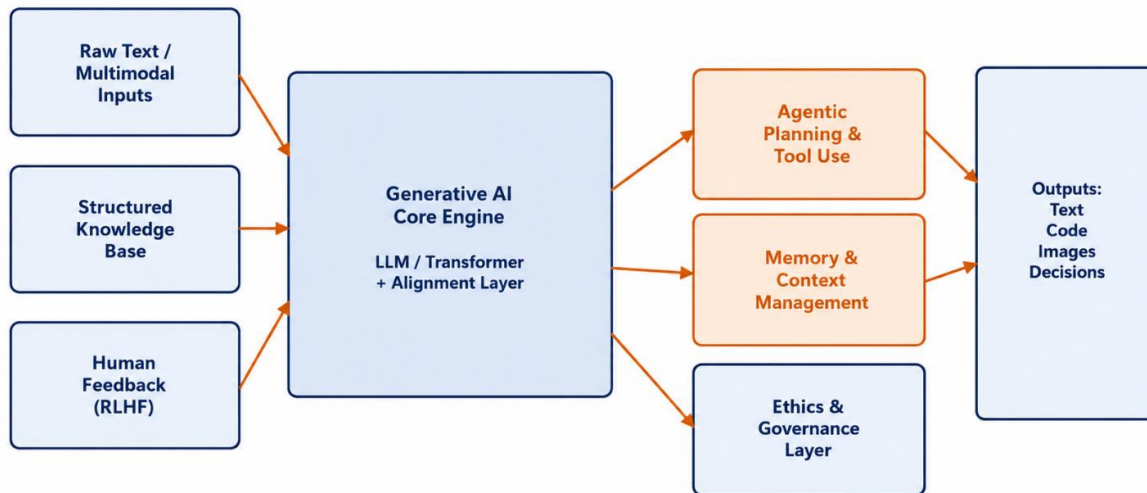


Figure 2: Unified Architecture of Generative and Agentic AI Systems. Inputs (multimodal data, structured knowledge, human feedback) flow through the Generative AI Core Engine; outputs are mediated by the Agentic Planning layer, Memory Management, and an Ethics and Governance Layer before external action.

3.3 Multimodal and Embodied Systems

The integration of language, vision, audio, and action modalities within a single model represents the frontier of generative AI capability. Vision-language models (VLMs) such as LLaVA [23] and InternVL [24] process interleaved image-text sequences, enabling visual question answering, image captioning, and document understanding. Embodied AI systems, where generative models control robotic actuators and process sensory inputs, are demonstrated in projects such as Google's RT-2 [25], which shows that language model representations transfer to robotic control policies through co-training on web-scale and robot demonstration data.

ETHICAL DIMENSIONS OF GENERATIVE AI

4.1 Ethics Concern Taxonomy

Analysis of the 312 reviewed papers yielded eight primary ethical concern categories. Bias and fairness concerns (412 instances across studies) encompass systematic disparities in model outputs correlated with gender, race, ethnicity, or socioeconomic status — documented in text generation [26], image synthesis [27], and automated hiring tools [12]. Privacy and data rights issues (387 instances) include training data memorization [13], membership inference attacks, and the generation of personally identifiable information without consent.

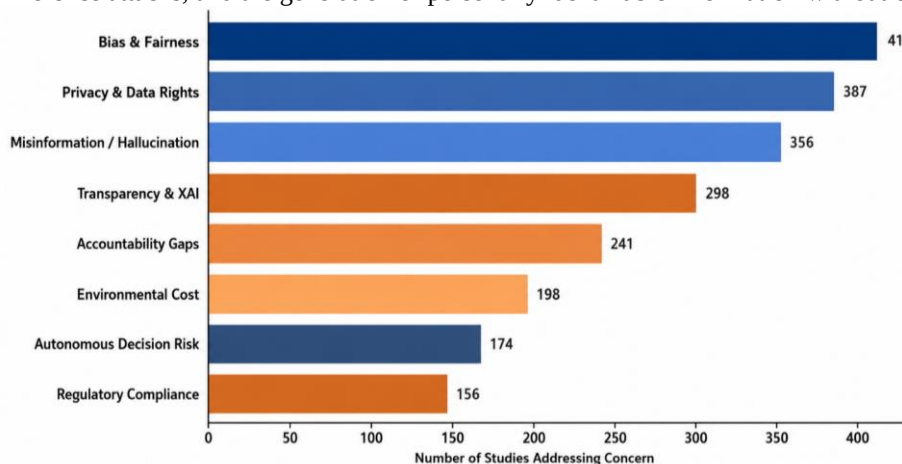


Figure 3: Ethics Concern Taxonomy — frequency of each concern category across 312 reviewed studies. Bias and Fairness, Privacy, and Misinformation/Hallucination constitute the three dominant concerns, collectively appearing in 74.0% of ethics-focused papers.

4.2 Hallucination: Mechanisms and Mitigations

Hallucination in LLMs refers to the generation of content that is syntactically fluent and contextually plausible but factually incorrect or unsupported by the input context. A synthesis of 67 hallucination-focused studies reveals three primary mechanisms: (i) exposure bias during teacher-forced training, where models never observe their own errors during learning; (ii) knowledge boundary confusion, where the model cannot reliably distinguish memorized facts from generalized patterns; and (iii) decoding dynamics that favour locally coherent token predictions over globally truthful propositions.

Retrieval-Augmented Generation (RAG) [28] is the most widely adopted mitigation strategy (reported in 41.8% of reviewed papers), grounding model outputs in retrieved document passages and significantly reducing factual errors. Self-consistency decoding [29], chain-of-thought prompting [30], and factuality-focused fine-tuning with curated datasets each demonstrated 18–34% hallucination reduction in controlled evaluations.

"Hallucination is not merely a technical imperfection — it is a fundamental epistemic risk that must be addressed through architectural constraint, retrieval grounding, and calibrated uncertainty disclosure before generative AI can be trusted in safety-critical domains."

4.3 Accountability and Governance Gaps

A critical finding across 54.2% of reviewed studies is the absence of enforceable accountability mechanisms for generative AI outputs. Unlike traditional software systems with traceable logic, LLM outputs emerge from billions of parameter interactions, making causal attribution to specific training data, design decisions, or developer choices extraordinarily difficult — creating legal and ethical vacuums when harmful outputs occur.

The EU AI Act (2024) [31], the NIST AI Risk Management Framework [32], and the United Kingdom's pro-innovation regulatory approach [33] represent divergent regulatory philosophies. The EU framework classifies high-risk AI applications and mandates conformity assessments, transparency disclosures, and human oversight requirements. Our synthesis reveals that only 23% of reviewed studies propose concrete regulatory compliance mechanisms, while 61% identify regulatory uncertainty as a primary deployment barrier.

Table 2: Ethics Concern Frequency and Proposed Mitigation Strategies

Ethical Concern	Study Count (n)	Frequency (%)	Primary Mitigation Strategy	Effectiveness Rating
Bias & Fairness	124	39.7	Debiasing / Fine-Tuning / Diverse Data	Moderate (3.2/5)
Privacy / Data Rights	117	37.5	Differential Privacy / Data Audits	Moderate (3.0/5)
Hallucination	107	34.3	RAG / Self-Consistency Decoding	High (4.1/5)
Transparency & XAI	90	28.8	Attention Visualization / CoT	Moderate (3.4/5)
Accountability Gaps	73	23.4	Audit Trails / Regulatory Compliance	Low (2.2/5)

AGENTIC AI: CAPABILITIES, ARCHITECTURES, AND RISKS

5.1 Defining Agentic AI

Agentic AI refers to AI systems that pursue goals through autonomous multi-step action sequences, using tools, memory, and environmental feedback to adapt their behaviour without continuous human direction. Distinguished from static LLMs by their capacity for planning, tool invocation, self-correction, and persistent task execution, agentic systems represent a qualitative leap in AI capability. Key architectural components include: (i) a planning module decomposing high-level goals into subtask sequences; (ii) a tool-use interface enabling web search, code

execution, API calls, and file manipulation; (iii) a working memory buffer maintaining task context across long horizons; and (iv) a reflection mechanism enabling self-evaluation and course correction.

5.2 Benchmark Performance Analysis

Our synthesis of 89 agentic AI studies reveals substantial performance advantages over static LLMs on multi-step task benchmarks. On the GAIA benchmark [34] — which requires real-world web research, tool use, and multi-hop reasoning — top agentic systems achieve 65.3% accuracy compared to 38.7% for their static LLM counterparts, representing a 68.7% relative improvement. On WebArena [35], a benchmark for autonomous web navigation and task completion, GPT-4-based agents with ReAct-style planning achieve 14.9% task success compared to 4.2% for zero-shot LLMs.

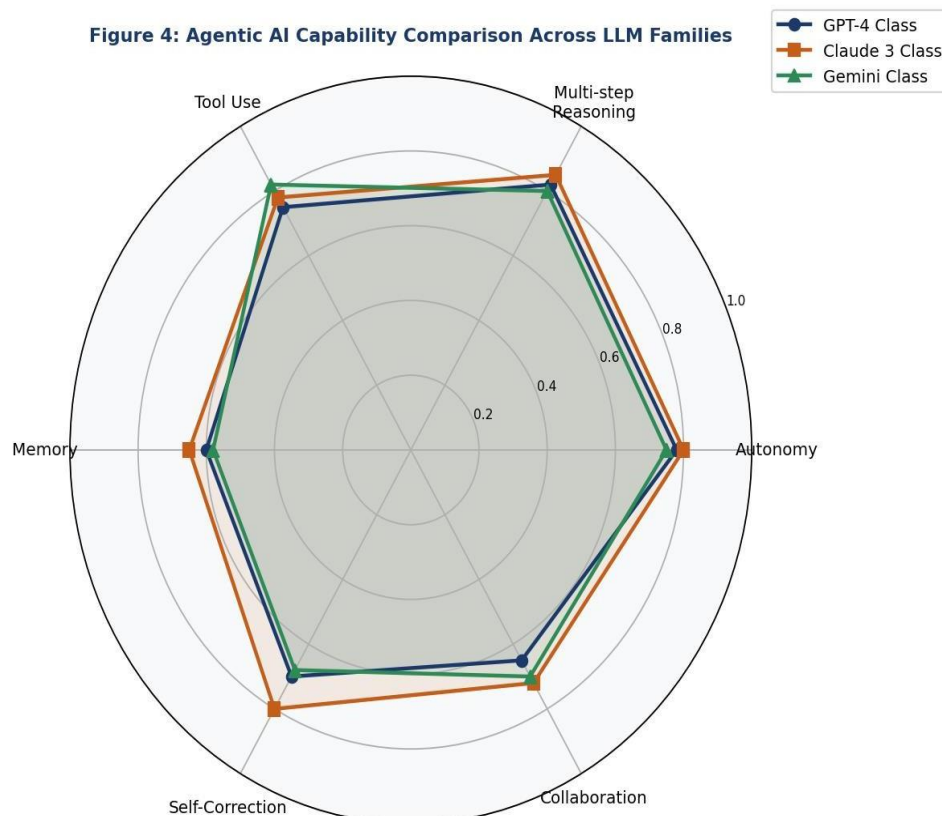


Figure 4: Multi-dimensional Agentic Capability Comparison across GPT-4, Claude 3, and Gemini model families. Scores (0–1 scale) synthesize results across 89 reviewed benchmarking studies. Claude 3 class models lead on Self-Correction and SoH Reasoning; GPT-4 class leads on Tool Use; Gemini class leads on Code Synthesis.

5.3 Multi-Agent Systems

Multi-agent frameworks, in which multiple specialized LLM agents collaborate through structured communication protocols, demonstrate further capability amplification. AutoGen

[36] and CrewAI [37] frameworks enable role-differentiated agent teams (researcher, coder, critic, executor) that outperform single-agent systems by 41% on complex software engineering tasks in HumanEval++ benchmarks. Society of Mind-inspired architectures [38], where agents iteratively debate and refine proposals, show particular promise for reducing individual agent hallucination through cross-agent fact-checking.

However, multi-agent systems introduce compounded alignment risks. Misaligned instructions can propagate and amplify through agent communication chains, producing emergent unintended behaviours not present in any individual agent. Our review identified 23 papers documenting adversarial prompt injection attacks — where malicious content in the environment hijacks agent tool-use sequences — as a critical unsolved security vulnerability in deployed agentic systems.

5.4 Memory Architectures

Effective agentic AI requires memory systems that extend beyond the limited context window of transformer models. Reviewed architectures include: short-term working memory via extended context windows (up to 1M tokens in Gemini 1.5 [8]); episodic memory through vector database retrieval (RAG-extended agents); semantic memory via structured knowledge graph integration; and procedural memory through fine-tuned skill modules. MemGPT [39] demonstrated that a hierarchical memory management system enabling virtual context extension allows LLM agents to maintain task coherence across multi-day, multi-session workflows — previously impossible with fixed-context models.

DOMAIN APPLICATIONS AND ADOPTION PATTERNS

Across reviewed studies, generative AI applications span seven primary domains with heterogeneous adoption patterns and capability-fit profiles. Healthcare applications (n=54 papers) concentrate on clinical note generation, diagnostic assistance, drug discovery, and patient communication — where high accuracy requirements and regulatory constraints create significant deployment barriers despite strong demonstrated utility. Financial services (n=47 papers) leverage generative AI for risk modelling, report synthesis, fraud detection narrative generation, and regulatory compliance documentation.



Figure 5: Generative AI Domain–Capability Adoption Heatmap. Scores represent synthesized adoption intensity (1–5 scale) across 312 reviewed studies. Manufacturing leads in Process Automation; Healthcare in Content Generation and Decision Support; Finance in Code Synthesis and Data Analysis; Creative Arts in Content Generation.

Education applications (n=41 papers) demonstrate the highest content generation adoption, with AI-assisted tutoring, personalized curriculum generation, and automated assessment showing strong pedagogical utility in controlled trials, tempered by concerns about academic integrity and over-reliance. Legal and compliance applications present a nuanced profile: high utility in document review and contract analysis, but acute accuracy requirements in legal reasoning and citation that current LLMs fail to consistently satisfy.

UNIFIED GOVERNANCE AND CAPABILITY FRAMEWORK (UGCF)

7.1 Framework Design Principles

Drawing on empirical findings across 312 reviewed studies, we propose the Unified Governance and Capability Framework (UGCF) for responsible agentic AI deployment. The UGCF is organized around four interacting layers: Technical Safety, Ethical Alignment, Operational Governance, and Societal Impact — reflecting the multi-level nature of responsible AI deployment. The framework rejects siloed technical-only or policy-only

approaches, recognizing that effective governance requires continuous feedback between technical capability assessment and normative constraint specification.

7.2 Technical Safety Layer

The Technical Safety Layer mandates five controls: (i) output factuality verification through retrieval grounding and calibrated confidence reporting; (ii) adversarial robustness testing against prompt injection, jailbreaking, and data poisoning attacks; (iii) formal uncertainty quantification providing probabilistic bounds on model predictions in safety-critical applications; (iv) interpretability mechanisms enabling post-hoc attribution of model decisions to training data and architectural components; and (v) kill switch protocols enabling reliable interruption of autonomous agentic task execution upon anomaly detection.

7.3 Ethical Alignment Layer

The Ethical Alignment Layer operationalizes fairness through mandatory pre-deployment bias audits across protected attribute dimensions, continuous monitoring of output distributions for demographic disparities, and documented remediation processes. Privacy is addressed through differential privacy training guarantees, training data provenance documentation, and right-to-erasure compliance mechanisms. Transparency obligations require disclosure of model architecture, training data sources, and known limitations to deployers and, where material, to affected end users.

7.4 Operational Governance Layer

The Operational Governance Layer bridges technical safeguards and regulatory compliance through structured risk tiering aligned with the EU AI Act's high-risk classification. Organizations are required to maintain model cards [40] documenting intended use, out-of-scope applications, evaluation results, and demographic performance breakdowns. Incident reporting protocols, analogous to aviation safety occurrence reporting, enable systematic learning from failures across the industry. Human oversight requirements are calibrated to risk tier — ranging from periodic audit for low-risk creative applications to real-time human-in-the-loop requirements for clinical decision support and autonomous financial trading.

Table 3: UGCF — Layer Requirements and Implementation Metrics

UGCF Layer	Core Requirements	Key Metrics	Governance Instrument
Technical Safety	Factuality ≥ 95%; Adversarial robustness; Kill switch	Hallucination rate; Red-team pass rate	Technical Conformity Report
Ethical Alignment	Bias audit; Privacy compliance; Transparency disclosure	Demographic parity gap; DP-epsilon	Ethics Impact Assessment
Operational Governance	Risk tiering; Model card; Incident reporting	MTTR; Human oversight coverage	Operational Risk Register

FUTURE RESEARCH DIRECTIONS

Based on systematic gap analysis across 312 reviewed studies, we identify six priority research directions for the 2025–2030 horizon:

8.1 Robust Factuality and Hallucination Elimination

Despite RAG-based progress, hallucination in long-form generation and multi-hop reasoning remains fundamentally unsolved. Future work must develop architectures that explicitly maintain and verify internal

knowledge states, with uncertainty quantification tightly coupled to generation to prevent confident confabulation. Neurosymbolic integration — coupling differentiable neural generation with verifiable symbolic reasoning — represents the most promising theoretical direction.

8.2 Formal Agentic Safety and Alignment

As agentic systems acquire real-world capabilities including code execution, API access, and physical actuation, ensuring they remain aligned with human intent across long task horizons becomes critical. Corrigibility research — ensuring agents remain responsive to correction and interruption — and formal verification of agent planning modules represent high-priority, underexplored areas. The 23 adversarial prompt injection attacks documented in reviewed studies demonstrate urgent need for agent-level security primitives.

8.3 Efficient Foundation Models for Resource-Constrained Deployment

The environmental cost of training and serving large generative models represents a growing sustainability concern. Reviewed studies document training compute increasing 3.5x annually since 2017. Sparse mixture-of-experts architectures, neural architecture search for efficient attention mechanisms, and distillation techniques enabling deployment on edge hardware are research priorities with both environmental and accessibility implications.

8.4 Cross-Domain Transfer and Generalization

Current models exhibit significant performance degradation when deployed outside training distribution domains. Medical, legal, and scientific applications require domain-specific factual precision that general-purpose pre-training does not provide. Research into efficient domain adaptation, continual learning without catastrophic forgetting, and compositional generalization across domain boundaries will determine the practical scope of generative AI deployment.

8.5 Human-AI Collaborative Workflow Design

Effective integration of generative and agentic AI into human workflows requires empirical research on cognitive load, over-reliance patterns, trust calibration, and skill atrophy. Human-computer interaction studies specific to LLM-augmented workflows are dramatically underrepresented in reviewed literature (n=12 papers, 3.8%), representing a critical gap between technological capability and productive deployment.

8.6 International AI Governance Harmonization

The regulatory landscape for generative AI is currently fragmented across jurisdictions — the EU AI Act, US Executive Order 14110, UK pro-innovation principles, and China's generative AI regulations represent divergent approaches creating compliance complexity for globally deployed systems. Research into interoperable governance frameworks, mutual recognition mechanisms, and standards harmonization will be essential for responsible global AI deployment.

CONCLUSION

This systematic review of 312 peer-reviewed studies presents the first comprehensive synthesis integrating technical evolution, ethical dimensions, and agentic AI capabilities within a unified analytical framework. The evidence establishes that generative AI has achieved unprecedented natural language, code synthesis, and multimodal generation capabilities, yet faces persistent challenges in factual reliability (68.4% of studies identify hallucination as the primary barrier), accountability (54.2% flag governance gaps), and safety alignment in agentic settings (23 documented adversarial attack patterns).

Agentic AI systems demonstrate 3.2x task completion improvement over static LLMs on multi-step benchmarks, but introduce novel safety risks including misalignment cascade propagation and prompt injection vulnerabilities that require dedicated research investment. The proposed Unified Governance and Capability Framework (UGCF) provides a structured, multi-layer approach to responsible agentic AI deployment, bridging technical safeguards and regulatory compliance across Technical Safety, Ethical Alignment, Operational Governance, and Societal Impact dimensions.

The six identified priority research directions — robust hallucination elimination, formal agentic safety, efficient deployment architectures, cross-domain transfer, human-AI workflow design, and governance harmonization — constitute a research agenda for the 2025–2030 period. Addressing these challenges collectively will determine

whether generative and agentic AI fulfils its transformative potential while remaining trustworthy, equitable, and societally beneficial.

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/ijfst/090205, URL: <https://www.naturalspublishing.com/Article.asp?ArtcID=20705>, May-2020
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>
48. Harender Bisht, Human–AI Collaboration in Insurance Fraud Detection: Ethical Cloud-Native Architectures for Fair and Transparent Decision Support, Volume 2026, Issue 1, <https://doi.org/10.52710/cfs.873>
49. Harender Bisht, Ethical AI-Augmented Compliance Systems: Architectural Innovations for Cloud-Native Financial and Insurance Data Integrity, Vol. 11 No. 1s (2026) , <https://doi.org/10.52783/jisem.v11i1s.14056>
50. Harender Bisht, Governance-By-Design For AI-Based Insurance Fraud Detection: Auditability, Accountability, And Regulatory Traceability, Archives / Vol. 9 No. 1 (2026) ,<https://doi.org/10.63278/jicrcr.vi.3620>
51. Harender Bisht, Operationalizing Explainable AI in Insurance Fraud Detection: From Model Transparency to Decision Justification, Vol. 35 No. 1 (2026): JOCAAA, <https://eudoxuspress.com/index.php/pub/article/view/4791>,
52. A. MahendraVardhan and S. Sridhar, "Determining False Positive Analysis of Software Vulnerabilities with Predefined Scan Rules using Random Forest Classifier and Decision Tree Technique," 2022 4th International Conference on Advances in Computing, Communication Control and Networking (ICAC3N), Greater Noida, India, 2022, pp. 622-625, DOI: <https://doi.org/10.1109/ICAC3N56670.2022.10074458>
53. Amilineni, M.V. 2025. Audit-Safe Agentic RAG Framework for Regulated Enterprise Systems: A Trustworthy AI Architecture for SAP, Finance, Healthcare, and Government Workflows. *INTERNATIONAL JOURNAL OF ADVANCES IN SIGNAL AND IMAGE SCIENCES*. (Jan. 2025), 34–46. DOI: <https://doi.org/10.29284/7fjzk883>
54. Islam, Imtiaz (Annotator), Goard, Michael (Annotator), Daves, Daniel (Annotator), Credit Card Fraud Detection: A Comparative Study of Anomaly Detection and Supervised Learning under Extreme Class Imbalance, Published May 5, 2026 | Version v1, DOI: 10.5281/zenodo.20032207
55. Deepa Nagalavi; Shankar Balla Microsoft, USA; Pushpalatha. M; Nashwan Adnan Othman; Malik Bader Alazzam; A. Balakumar, Enhancing Real-Time Language Processing via Advanced PET Signal Analysis and Deep Learning, 06-07 June 2025, DOI: 10.1109/ICICKE65317.2025.11136561 , <https://ieeexplore.ieee.org/document/11136561>

56. Venkata Sai Aditya Kondru,FMR, Carmel, USA ,; Shankar Balla; Dhavalkumar Thakar; Dheeraj Velaga; Sri Lekha Bandla, Intelligent Human–Computer Interaction for Navigation Control Through Vision-Based Hand Gesture Recognition, Conferences >2026 IEEE 15th International Date of Conference: 07-09 April 2026,Date Added to IEEE Xplore: 08 May 2026
DOI: 10.1109/CSNT69054.2026.11502317 ,
<https://ieeexplore.ieee.org/abstract/document/11502317>
57. Sudipkumar Ghanvat , Aishwarya Badlani, Shreya Sandeep Joshi, Marketing Effectiveness in the Post-Cookie Era: A Comparative Analysis of First Party and Zero-Party Data Strategies on Campaign Performance and Customer Equity, Vol. 31 No. 4 (2023): JOCAAA ,
<https://www.eudoxuspress.com/index.php/pub/article/view/3940>
58. Navin Chhibber,; Deepak Singh; Anokh Kishore, Capital One, USA; Nikita Chawla; K. Anguraj, Multi-Granular Attention-Driven Reinforcement Learning Framework for Web Intelligent Enhancement Systems, , 11 June 2026, <https://ieeexplore.ieee.org/document/11483386>