

ORIGINAL RESEARCH

Open access

A Theory of Reproducibility for AI-Generated Materials Insights: What Must Be True for Claims to Hold

Wei Liu^{1*}, Zhang Min¹

Abstract

The integration of artificial intelligence (AI) into materials science has transformed the landscape of discovery and insight generation, enabling rapid analysis of complex datasets and simulation of material behaviors at unprecedented scales. However, the reproducibility of AI-generated insights remains a pivotal concern, as it underpins the epistemic validity of claims derived from such systems. This conceptual paper develops a novel theoretical framework that interprets reproducibility not as a static attribute but as an emergent property arising from dynamic interactions among data ecosystems, algorithmic architectures, and human interpretive practices. By synthesizing literature on AI trustworthiness and materials informatics, the framework elucidates the systemic conditions—such as data lineage transparency, algorithmic feedback loops, and ethical epistemic alignments—that must align for AI-derived claims to sustain scrutiny across contexts. It emphasizes interaction dynamics where data quality influences model robustness, while human oversight modulates algorithmic outputs, fostering a balanced ecosystem for reliable insights. Ethical reasoning is integrated throughout, highlighting trade-offs between computational efficiency and interpretive depth. This approach shifts focus from isolated reproducibility metrics to holistic systems-level insights, offering guidance for scholars and practitioners in applied AI for materials science. Ultimately, the framework advocates for a steering logic that prioritizes integrative processes over predictive assertions, ensuring that AI contributions enhance rather than undermine the foundational integrity of materials knowledge.

Keywords Conceptual framework, AI reproducibility, Materials insights, Data ecosystems, Algorithmic dynamics, Epistemic ethics

*Correspondence:

Wei Liu

wei.liu@outlook.com

¹ Department of Computational Materials Science, School of Materials Engineering, Shanghai Jiao Tong University, Shanghai, China

Introduction

The advent of artificial intelligence (AI) in materials science marks a paradigm shift in how researchers conceptualize and pursue knowledge about material properties, structures, and behaviors. AI systems, particularly those leveraging machine learning algorithms, process vast amounts of data to uncover patterns that might elude traditional analytical methods, thereby accelerating discovery in fields such as nanomaterials, alloys, and polymers [1, 2]. This acceleration is evident in the way AI facilitates the exploration of high-dimensional parameter spaces, where conventional experimentation would be prohibitively time-consuming or resource-intensive [3]. Yet, as AI becomes increasingly embedded in the generation of materials insights, questions about the reproducibility of these insights emerge as central to maintaining the discipline's epistemic rigor. Reproducibility, in this context, extends beyond mere replication of results; it encompasses the conditions under which AI-generated claims can be interpreted as holding across diverse settings, users, and temporal scales.

At its core, reproducibility in AI-generated materials insights involves a multifaceted interplay of technical, methodological, and philosophical elements. Technically, it requires transparency in data handling and algorithmic operations, ensuring that the pathways from raw inputs to derived conclusions are traceable and verifiable [4]. Methodologically, it demands consistency in how models are trained, validated, and applied, accounting for variations in computational environments and data sources [5]. Philosophically, it invokes epistemic considerations about what constitutes valid knowledge in a domain where human intuition intersects with machine computation [6]. In materials science, where insights often inform real-world applications such as sustainable energy storage and advanced composites, the stakes are high: irreproducible claims could lead to misguided investments or flawed innovations [7].

The urgency of addressing reproducibility stems from the inherent complexities of AI systems. Unlike deterministic simulations rooted in physical laws, AI models in materials science often operate as black-box entities, with their internal decision-making processes opaque [8]. This opacity can obscure the dependencies between input data quality and output reliability, making it challenging to discern whether an insight reflects genuine material phenomena or artifacts of the model's training regime [9]. Moreover, the rapid evolution of AI techniques—such as deep learning architectures tailored for molecular property prediction—introduces variability, further complicating reproducibility [10]. For instance, subtle differences in hyperparameter tuning or dataset preprocessing can yield divergent outcomes, even when applied to the same problem domain [11].

Literature highlights several systemic challenges that undermine reproducibility. Data heterogeneity, arising from disparate experimental protocols or simulation standards, often introduces inconsistencies that propagate through AI pipelines [12]. Algorithmic brittleness, where models perform well on training data but falter in novel contexts, exacerbates this issue [13]. Additionally, the human element—researchers' interpretive biases in selecting features or evaluating results—adds another layer of variability [14]. These challenges are not isolated; they interact dynamically, creating feedback loops that can amplify uncertainties in AI-generated insights [15].

To navigate these complexities, a conceptual approach is essential, one that interprets reproducibility as an emergent outcome of integrated systems rather than a checklist of isolated criteria. This paper posits that for AI-generated claims in materials science to hold, certain systemic alignments must prevail: data must be contextualized within robust ecosystems, algorithms must incorporate adaptive logics, and human-AI interactions must foster ethical epistemic reasoning [16]. Such alignments enable trade-offs, like balancing model complexity with interpretability, to be managed thoughtfully [17].

The theoretical foundation draws from interdisciplinary insights, blending concepts from informatics, philosophy of science, and systems theory. For example, systems-level perspectives emphasize how components like data repositories and computational frameworks interconnect to produce stable knowledge [18]. Ethical reasoning underscores the responsibility to ensure that AI insights align with broader scientific values, such as inclusivity and sustainability [19]. By focusing on these dynamics, the paper avoids reductive formulations and instead explores how reproducibility emerges from ongoing interactions.

This conceptual exploration is timely, as AI adoption in materials science surges amid calls for trustworthy computing [20]. Recent advancements, such as AI-driven high-throughput screening for catalysts, demonstrate potential but also reveal reproducibility gaps when models are transferred across laboratories [21]. Addressing these gaps requires a framework that interprets the conditions for claim validity through lenses of interaction and integration, rather than static assertions.

In synthesizing existing literature, the paper identifies key themes: the role of data in grounding AI outputs, the dynamics of algorithmic decision-making, and the epistemic trade-offs in human-AI collaboration [22]. These themes inform the proposed framework, which conceptualizes reproducibility as a steering mechanism guiding the evolution of materials knowledge [23].

Ultimately, this work contributes to applied AI in materials science by offering a novel interpretive lens. It encourages scholars to view reproducibility not as a hurdle but as a foundational dynamic that enhances the discipline's resilience. By delineating what must align for claims to endure, the framework provides a pathway to more robust, ethically grounded insights, fostering a future in which AI amplifies human ingenuity without compromising epistemic integrity [24].

Theoretical Background and Literature Synthesis

Evolution of AI in materials science

The application of AI in materials science has evolved from auxiliary tools for data analysis to central drivers of insight generation, reshaping how researchers interpret material behaviors and properties. Early integrations focused on machine learning for pattern recognition in experimental datasets, enabling faster identification of correlations in crystal structures or alloy compositions [1, 25]. Over the past few years, advancements in deep learning have expanded this scope, enabling the simulation of quantum-level interactions and the prediction of emergent properties in complex materials [2, 26]. This evolution reflects a shift toward data-driven paradigms, where AI systems process multimodal inputs—combining experimental, computational, and literature-derived data—to yield insights that inform design strategies [3, 27].

Central to this progression is the concept of materials informatics, which interprets AI as a bridge between raw data and conceptual understanding [4]. For instance, AI enables navigation through vast chemical spaces, where traditional methods falter due to combinatorial explosion [5]. However, this reliance on AI introduces interpretive challenges, as the generated insights must be contextualized within physical principles to avoid misattribution of causality [6, 28]. Literature underscores the dynamic interplay between AI's computational power and the domain-specific knowledge it augments, highlighting how such systems enable integrative reasoning across scales—from atomic to macroscopic [7].

Challenges in the reproducibility of AI outputs

Reproducibility in AI-generated materials insights encounters systemic hurdles stemming from the intricate dynamics of data processing and model execution. Data variability, including inconsistencies in measurement standards or simulation parameters, often leads to divergent interpretations of the same material phenomena [8, 29]. This variability creates feedback structures in which initial data discrepancies amplify across algorithmic layers, resulting in outputs that lack consistency across replications [9].

Algorithmic opacity further complicates these dynamics, as neural networks' internal logic remains largely inscrutable, hindering the tracing of how inputs translate into claims [10, 30]. Systems-level insights reveal that such opacity fosters epistemic uncertainty, in which users struggle to discern whether insights arise from genuine patterns or model artifacts [11]. Ethical considerations emerge here, as irreproducible insights could perpetuate inequities in resource allocation for materials research [12, 31].

Moreover, human-AI interactions introduce additional layers, in which interpretive practices vary by researchers' expertise, leading to inconsistent claim validation [13]. Trade-offs abound, such as between model accuracy and computational reproducibility, where resource constraints in different settings yield varying outcomes [14, 32].

Frameworks for trustworthiness in AI

Existing frameworks for AI trustworthiness provide interpretive foundations for addressing reproducibility, emphasizing transparency and robustness as key dynamics. Concepts like generalizability and explainability interpret AI systems as needing to balance predictive power with verifiable logics [15, 33]. In materials science, these frameworks advocate for data curation practices that ensure lineage traceability, fostering ecosystems where insights can be cross-verified [16].

Interaction dynamics between components—such as data sources, models, and validation protocols—form the core of these approaches, highlighting feedback loops that enhance stability [17, 34]. Ethical epistemic reasoning integrates, urging alignments that prioritize fairness in how AI insights are generated and applied [18]. Systems-level insights from these frameworks reveal trade-offs, like sacrificing some model flexibility for greater reproducibility [19, 35].

Data ecosystems and their role in insight generation

Data ecosystems in materials science serve as the bedrock for AI reproducibility, interpreting quality not as isolated attributes but as relational dynamics within broader networks. High-quality data, characterized by completeness and consistency, enables robust model training, but literature notes that fragmented repositories often disrupt these dynamics [20, 21]. Integrative processes, such as standardized ontologies, facilitate smoother interactions between disparate data sources, enhancing the flow of information [22].

Epistemic reasoning underscores the need for ethical stewardship in data handling, where biases in sourcing can skew AI outputs [23]. Systems-level views interpret these ecosystems as adaptive structures, where feedback from usage refines data over time [24, 25].

Algorithmic logics and feedback structures

Algorithmic logics in AI for materials insights involve complex steering mechanisms that guide data toward meaningful outputs. Deep learning architectures, for example, embed hierarchical feedback loops that adapt to input variations, but this adaptability can introduce instabilities that affect reproducibility [26, 27]. Conceptual interpretations frame these logics as balancing exploration of novel patterns with adherence to known physical constraints [28].

Trade-offs emerge in designing these structures, such as between innovation and reliability, where overly rigid logics may miss insights, while flexible ones risk irreproducibility [29, 30]. Ethical dimensions call for logics that incorporate accountability, ensuring outputs align with scientific norms [31].

Human-AI interaction dynamics

Human-AI interactions represent a critical nexus for reproducibility, where interpretive practices shape how claims are formulated and validated. Systems-level insights highlight collaborative dynamics, with humans providing contextual oversight that modulates AI outputs [32, 33]. Feedback structures in these interactions enable iterative refinement, thereby enhancing claim robustness [34].

Ethical epistemic reasoning emphasizes equitable partnerships, avoiding over-reliance on AI that could erode human expertise [35]. Trade-offs include balancing automation efficiency with interpretive depth, ensuring that claims hold through integrated human-machine reasoning [1, 2].

Proposed conceptual framework

The proposed framework conceptualizes reproducibility in AI-generated materials insights as an emergent epistemic property arising from sustained alignment among data ecosystems, algorithmic architectures, and human interpretive practices. Reproducibility is not treated as a fixed criterion or an outcome that can be secured through isolated controls; instead, it is interpreted as a dynamic condition that depends on the coherence of interactions across these domains. Claims derived from AI systems are understood to hold only insofar as these interactions remain epistemically aligned over time. At the center of the framework lies a steering logic that governs how trade-offs—such as those between computational scalability and interpretive depth—are navigated to preserve claim validity.

Data ecosystems constitute the foundational domain of the framework and are interpreted as adaptive epistemic environments rather than passive repositories. Within this view, data quality emerges relationally through provenance,

contextual integrity, and interoperability across experimental, simulated, and archival sources. Feedback structures within the ecosystem continuously reshape the epistemic grounding of AI-generated claims, as new data exposures refine or destabilize prior interpretations. This dynamic foregrounds ethical epistemic reasoning, insofar as inclusive curation and contextual transparency are necessary to prevent the entrenchment of bias and to sustain the intelligibility of claims across settings.

Algorithmic architectures operate as mediating interpretive engines layered upon these data ecosystems. Rather than functioning as neutral processors, algorithms are conceptualized as dynamic structures whose internal logics shape how data is rendered into claims. Adaptive feedback mechanisms allow algorithms to respond to contextual variation, but this adaptability introduces epistemic tension between exploratory flexibility and interpretive stability. Systems-level insights arise from how algorithmic behavior aligns—or fails to align—with the contextual assumptions embedded in data ecosystems. When such alignments are maintained, reproducibility is reinforced; when they fracture, epistemic fragilities emerge that compromise the endurance of claims.

Human interpretive practices complete the triadic structure of the framework and occupy a constitutive, rather than supervisory, role. Human judgment, contextual reasoning, and disciplinary norms are not external correctives applied after computation, but integral components that stabilize meaning and legitimacy. Through iterative interaction with algorithmic outputs, human agents modulate interpretive boundaries, negotiating trade-offs between transparency and opacity, and between automation and epistemic scrutiny. These interactions generate closed-loop feedback through which interpretations reshape data practices and algorithmic assumptions, reinforcing systemic coherence.

The integrative force of the framework lies in its treatment of ethical epistemic reasoning as a binding condition across all domains. Reproducibility is interpreted as arising from harmonious alignment rather than control, such that disturbances in one component propagate through the system and require compensatory realignment elsewhere. For AI-generated materials to hold claims, data, algorithms, and interpretation must remain mutually intelligible and normatively aligned. The framework thus characterizes reproducibility as a structural property of epistemic ecosystems, enabling a resilient mode of knowledge production that sustains AI-generated insights in materials science over time. The epistemic roles of the framework’s core components, along with the conditions under which misalignment destabilizes claim validity, are synthesized in [Table 1](#).

Table 1. Epistemic components and alignment conditions for reproducibility in AI-generated materials insights

Framework component	Epistemic role	Source of stability	Typical misalignment	Consequence for claim validity
Data ecosystems	Ground material claims in contextualized evidence	Provenance, contextual integrity, interoperability	Fragmented datasets, opaque lineage, context loss	Claims lose intelligibility across settings and cannot be meaningfully contested
Algorithmic architectures	Transform data into inferential structures	Alignment between internal logics and data assumptions	Opacity, over-adaptivity, brittle generalization	Claims appear reproducible locally but fail under contextual transfer
Human interpretive practices	Stabilize meaning and epistemic legitimacy	Domain knowledge, judgment, normative reasoning	Automation bias, interpretive disengagement	Claims lack epistemic warrant despite computational consistency
Ethical epistemic reasoning	Bind system components into legitimate knowledge production	Normative alignment, accountability, transparency	Unequal access, hidden value assumptions	Reproducibility privileges actors with infrastructural power

Steering logic (central nexus)	Navigate trade-offs and sustain alignment	Continuous feedback and interpretive coherence	Static optimization or isolated control	Systemic fragility despite local performance
--------------------------------	---	--	---	--

Figure 1 illustrates the tripartite framework linking data ecosystems, algorithmic architectures, and human interpretive practices.

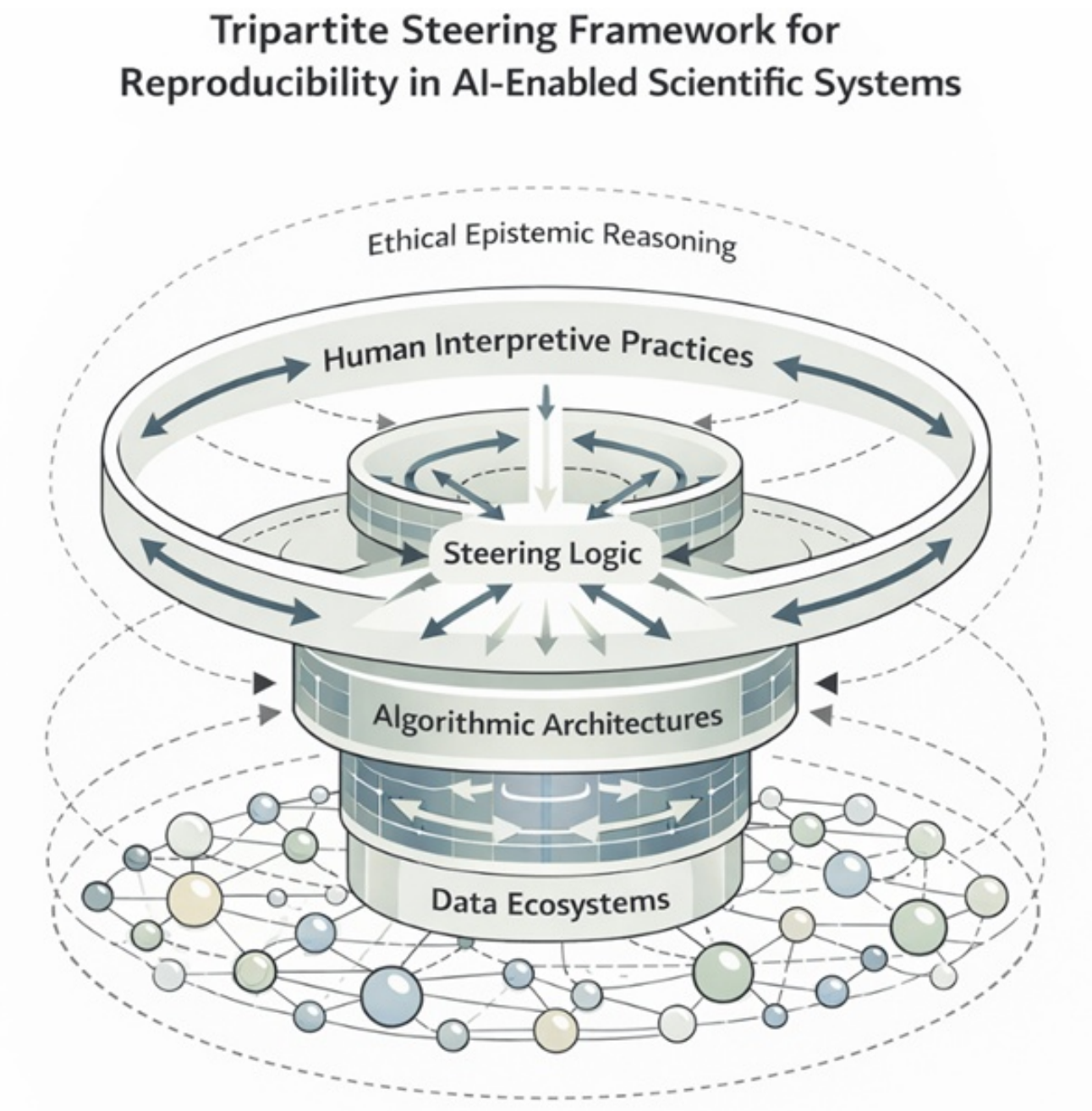


Figure 1. Tripartite framework depicting dynamic interactions among data, algorithms, and human interpretation in reproducible materials-AI systems

Analytical implications

The conceptual framework advanced in this paper repositions reproducibility in AI-generated materials insights as an epistemic condition emerging from systemic alignment, rather than as a technical property of models or datasets. Analytically, this reframing alters how the reliability of AI outputs is interpreted: claims are no longer evaluated as isolated artifacts, but as knowledge statements whose validity depends on the coherence of interacting data, algorithmic, and interpretive structures over time. Reproducibility, in this sense, functions as a measure of epistemic continuity rather than procedural repeatability. The principal epistemic trade-offs and feedback structures that stabilize or undermine reproducibility are summarized in **Table 2**.

Table 2. Epistemic trade-offs and feedback structures shaping reproducibility in materials AI

Epistemic tension	Competing poles	Feedback mechanism	Effect on reproducibility
Scalability vs. interpretive depth	High-throughput automation ↔ Contextual reasoning	Human–AI interpretive loops	Excessive automation weakens claim endurance without interpretive reinforcement
Algorithmic flexibility vs. stability	Adaptive learning ↔ Inferential consistency	Internal model feedbacks	Over-adaptivity destabilizes claims across contexts
Data breadth vs. contextual fidelity	Large heterogeneous datasets ↔ Local meaning	Data lineage refinement	Broad data without context dilutes epistemic grounding
Transparency vs. performance	Explainability ↔ Predictive optimization	Ethical epistemic steering	Opaque performance undermines trust in claim validity
Efficiency vs. epistemic accountability	Speed of insight ↔ Justification rigor	Normative oversight	Fast insights without accountability erode legitimacy

From the perspective of data ecosystems, the framework implies that reproducibility is analytically grounded in the traceability and contextual integrity of informational flows. Data quality is not treated as an intrinsic attribute, but as a relational property shaped by provenance, curation practices, and interoperability across sources. Failures of reproducibility are therefore interpretable as breakdowns in epistemic lineage, in which misaligned data contexts undermine the conditions for claims to remain intelligible and contestable across settings. This interpretation foregrounds trade-offs between breadth and epistemic depth, highlighting how expansive data aggregation may weaken claim stability if contextual grounding is not preserved.

At the level of algorithmic architectures, the framework introduces an analytical distinction between computational adaptability and epistemic stability. Reproducibility emerges not from algorithmic rigidity, but from the alignment between internal model dynamics and the interpretive structures that render outputs meaningful. Model brittleness and opacity are thus understood as epistemic risks rather than purely technical shortcomings, insofar as they obscure the inferential pathways through which claims are generated. From this viewpoint, algorithmic feedback structures shape reproducibility by either stabilizing or destabilizing the interpretive coherence of outputs across variations in data and context.

Human interpretive practices occupy a central analytical role within the framework, not as external validators, but as constitutive elements of reproducibility itself. The framework implies that AI-generated insights achieve epistemic standing only through sustained interpretive alignment with domain knowledge, conceptual expectations, and normative standards of reasoning within materials science. Reproducibility is therefore co-constructed through human–AI interaction, where judgment, contextualization, and critical scrutiny act as stabilizing forces. Trade-offs emerge between efficiency and interpretive depth, revealing how accelerated automation may erode epistemic robustness if human engagement is reduced to procedural oversight rather than substantive reasoning.

At the systems level, the analytical implications extend beyond individual components to the structure of materials knowledge production itself. By interpreting reproducibility as an emergent property, the framework suggests that AI can contribute durable insights only when systemic coherence is maintained across data practices, algorithmic behavior, and interpretive norms. Disruptions in any one domain propagate through the system, necessitating compensatory realignments elsewhere. Reproducibility failures thus signal systemic epistemic fragility rather than isolated error, reframing how responsibility for AI-generated claims is distributed across infrastructures and practices.

Ethically, these analytical implications situate reproducibility as a condition of epistemic legitimacy rather than mere technical adequacy. Claims that cannot be sustained across contexts risk consolidating asymmetric epistemic authority, privileging actors with the resources to stabilize fragile AI pipelines. From this perspective, reproducibility encompasses questions of fairness, accountability, and the equitable distribution of epistemic credibility within the field. Ethical epistemic reasoning, therefore, operates not as an external constraint but as an internal requirement for claims to warrant trust.

Taken together, these analytical implications encourage a reorientation of epistemic practices in AI-driven materials science. Reproducibility emerges as a steering principle that governs how AI-generated insights are interpreted, stabilized, and integrated into the evolving body of materials knowledge. Rather than constraining innovation, this perspective clarifies the conditions under which AI can function as a legitimate contributor to scientific understanding, ensuring that acceleration of discovery does not come at the expense of epistemic integrity.

Results and Discussion

By conceptualizing reproducibility as an emergent property arising from dynamic interactions, the proposed framework advances current discussions of AI trustworthiness in materials science beyond fragmented, component-level treatments. Much of the existing literature on data ecosystems aligns with the view that data quality is relational rather than intrinsic, emphasizing how inconsistencies propagate through AI pipelines when contextual alignment is absent [1, 2]. The present framework extends these accounts by situating data lineage transparency within a broader epistemic structure, where reproducibility depends not only on technical traceability but also on ethical judgments regarding data sharing, openness, and stewardship—dimensions that are often acknowledged but insufficiently theorized in prior work [35].

At the algorithmic level, the framework aligns with established discussions on robustness and generalization, where adaptive feedback mechanisms are considered essential for maintaining performance across variable contexts [4, 5]. However, the framework contributes a distinct interpretive advance by reframing algorithmic opacity as a systemic epistemic disturbance rather than a localized technical limitation. From this perspective, opacity undermines reproducibility not merely by obscuring internal logic, but by disrupting the alignment between computational behavior and the interpretive structures required to stabilize claims. This interpretation complements emerging work in explainable AI [6, 34] while extending it toward a systems-level understanding of how trade-offs between complexity and interpretability shape the endurance of AI-generated knowledge claims [8, 9].

Human–AI interaction occupies a more constitutive role in the proposed framework than in much of the existing literature on collaborative or human-in-the-loop systems [10, 11]. Whereas prior studies often treat human oversight as an external corrective or validation step, the framework interprets human interpretive practices as integral to the epistemic conditions of reproducibility itself. This co-constructive dynamic implies a process of mutual adaptation, in which human judgment and algorithmic outputs evolve together, shaping what counts as a stable and credible claim. In doing so, the framework foregrounds epistemic ethics as a stabilizing force, particularly in mitigating interpretive biases that can otherwise be amplified through automated reasoning [12, 13]. Systems-level insights further distinguish this approach from studies that isolate human oversight, demonstrating how interpretive engagement contributes to resilience across the entire knowledge-production system [14, 32].

Within broader materials informatics discourse, reproducibility is increasingly recognized as a systems property rather than a procedural criterion [16, 33]. The present framework reinforces this shift while adding analytical depth by explicitly

modeling interaction dynamics and feedback loops across scales. In contrast to accounts that emphasize standardization or benchmarking in isolation, the framework highlights how reproducibility depends on continuous alignment among evolving data practices, algorithmic behaviors, and disciplinary norms of interpretation [18, 19]. Ethical reasoning is woven throughout this analysis, extending prevailing responsible AI narratives by explicitly addressing trade-offs in resource allocation, access, and the uneven capacity of research communities to sustain reproducible AI pipelines [20, 21].

Persistent challenges in the literature—such as data scarcity, heterogeneity, and uneven infrastructural support—resonate particularly with the framework's conception of adaptive data ecosystems [22, 23]. Rather than treating these challenges as deficits to be corrected through isolated technical interventions, the framework interprets them as structural conditions that shape epistemic vulnerability. In materials domains characterized by high complexity and contextual sensitivity, misalignments across system components are shown to generate epistemic fragility, in which claims may appear robust locally yet fail to endure across settings or over time [24, 30].

The central contribution of the framework lies in its interpretive synthesis. Moving beyond reductive reproducibility metrics, it offers a nuanced account of what sustains AI-generated claims as legitimate contributors to materials knowledge [26, 27]. This synthesis engages productively with interdisciplinary perspectives from informatics, philosophy of science, and systems theory, enriching ongoing debates about AI's role in accelerating discovery while preserving epistemic grounding [28, 31]. In doing so, the framework positions reproducibility not as a constraint on innovation, but as a structural condition that enables AI to function as a durable and trustworthy epistemic instrument within materials science.

Conclusion

This paper advances a theory of reproducibility for AI-generated materials insights by reframing reproducibility as an epistemic condition that emerges from sustained alignment across data ecosystems, algorithmic architectures, and human interpretive practices. Rather than treating reproducibility as a procedural requirement or a post hoc verification step, the framework establishes it as a structural property of knowledge systems—one that determines whether AI-generated claims can justifiably endure across contexts, users, and time.

By articulating reproducibility as a systems-level phenomenon, the framework clarifies how interaction dynamics and feedback structures govern the stability of AI-derived claims. Data lineage, algorithmic behavior, and human judgment are shown to be inseparable in shaping epistemic continuity, such that disruptions in any single component propagate across the system. Ethical epistemic reasoning functions not as an external constraint but as an internal organizing principle, guiding how trade-offs between innovation, interpretability, and accountability are navigated in the production of materials knowledge.

The central contribution of this work lies in its interpretive integration. It provides a coherent lens through which reproducibility is understood not as a limiting factor for AI-driven discovery, but as the condition that enables AI systems to function as legitimate epistemic instruments within materials science. By specifying what must remain aligned for claims to hold, the framework offers a durable conceptual foundation for evaluating AI-generated insights without reducing scientific validity to mere technical replication.

As applied artificial intelligence continues to reshape materials research, this theory of reproducibility reorients scholarly attention toward the epistemic infrastructures that sustain knowledge over time. In doing so, it positions reproducibility as a steering principle for the responsible evolution of AI in materials science—ensuring that acceleration of insight remains anchored in epistemic integrity.

Acknowledgements

None

Conflict of interest

None

Financial support

None

Ethics statement

None

Received: 06 Mar 2024

Revised: 20 Apr 2024

Accepted: 02 Jun 2024

Published online: 18 July 2024

Rights and permissions

Open Access This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

References

- Suh C, Fare C, Warren JA, Pyzer-Knapp EO. Evolving the materials genome: How machine learning is fueling the next generation of materials discovery. *Annu Rev Mater Res.* 2020;50:1-25.
- Batra R, Song L, Ramprasad R. Emerging materials intelligence ecosystems propelled by machine learning. *Nat Rev Mater.* 2021;6:655-78.
- Liu L, Bi M, Wang Y, Liu J, Jiang X, Xu Z, et al. Artificial intelligence-powered microfluidics for nanomedicine and materials synthesis. *Nanoscale.* 2021;13:19352-66.
- Merchant A, Batzner S, Schoenholz SS, Aykol M, Cheon G, Cubuk ED. Scaling deep learning for materials discovery. *Nature.* 2023;624:80-5.
- Szymanski NJ, Bartel CJ, Zeng Y, Tu Q, Corman M, et al. An autonomous laboratory for the accelerated synthesis of novel materials. *Nature.* 2023;624:86-91.
- Pyzer-Knapp EO, Chen L, Day GM, Cooper AI. Accelerating computational discovery of porous solids through improved navigation of energy-structure-function maps. *Sci Adv.* 2021;7:eabi4763.

Sha W, Guo Y, Yuan Q, Tang S, Zhang X, Lu S, et al. Artificial Intelligence to Power the Future of Materials Science and Engineering. *Adv Intell Syst.* 2020;2:1900143.

Senior AW, Evans R, Jumper J, Kirkpatrick J, Sifre L, Green T, et al. Improved protein structure prediction using potentials from deep learning. *Nature.* 2020;577:706-10.

Jumper J, Evans R, Pritzel A, Green T, Figurnov M, Ronneberger O, et al. Highly accurate protein structure prediction with AlphaFold. *Nature.* 2021;596:583-9.

Morgan D, Jacobs R. Opportunities and challenges for machine learning in materials science. *Annu Rev Mater Res.* 2020;50:71-103.

Kwaria RJ, Mondarte EAQ, Tahara H, Chang R, Hayashi T. Data-driven prediction of protein adsorption on self-assembled monolayers toward material screening and design. *ACS Biomater Sci Eng.* 2020;6:4949-56.

Goswami A. A Review on Applications of Artificial Intelligence in Material Engineering. *Adv Eng Mater.* 2023;25:2300104.

Zhou Y, Ping X, Guo Y, Heng BC, Wang Y, Meng Y, et al. Assessing Biomaterial-Induced Stem Cell Lineage Fate by Machine Learning-Based Artificial Intelligence. *Adv Mater.* 2023;35:2210637.

Al-Kharusi G, Dunne NJ, Little S, Levingstone TJ. The role of machine learning and design of experiments in the advancement of biomaterial and tissue engineering research. *Bioengineering.* 2022;9:561.

Fang J, Wang J, Li Y, Song Y, Wang Q, Sun M, et al. Machine learning accelerates the materials discovery. *Mater Today Sustain.* 2022;18:100143.

Cai J, Chu X, Wei Y, Yang J, Zhang C, Gao H, et al. Machine learning-driven new material discovery. *Nanoscale Adv.* 2020;2:3115-22.

Li KQ, O'Farrell H, Oliver JA, Yang L. Estimating the thermal conductivity of soils using six machine learning algorithms. *Int Commun Heat Mass Transfer.* 2022;136:106149.

Walker E, Paniagua S, Chugh T, Le Bras R, Balasubramanian G, Halloran B, et al. Recent advances and applications of deep learning methods in materials science. *npj Comput Mater.* 2022;8:125.

Choudhary K, Bercx M, Tavazza F. Machine learning with force-field-inspired descriptors for materials: Fast screening and mapping energy space. *Phys Rev Mater.* 2021;5:083801.

Dan Y, Zhao Y, Li X, Li S, Hu M, Hu J. Generative adversarial networks (GAN) based efficient sampling of chemical composition space for inverse design of inorganic materials. *npj Comput Mater.* 2020;6:84.

Saidi WA, Shadid W, Vesely EJ. Cross-domain multiscale data integration for structure-property predictions of organic crystals. *Chem Mater.* 2020;32:4882-91.

Boztepe C, Yüceer M, Inan T. Prediction of the deswelling behaviors of pH- and temperature-responsive poly(NIPAAm-co-AAc) IPN hydrogel by artificial intelligence techniques. *Res Chem Intermed.* 2020;46:409-28.

Fung V, Zhang J, Juarez E, Sumpter BG. Benchmarking graph neural networks for materials chemistry. *npj Comput Mater.* 2021;7:84.

Chen L, Tran H, Batra R, Kim C, Ramprasad R. Machine learning models for the prediction of energy, forces, and stresses for platinum. *npj Comput Mater.* 2021;7:19.

Chen C, Zuo Y, Ye W, Li X, Ong SP. Learning properties of ordered and disordered materials from multi-fidelity data. *Nat Comput Sci.* 2021;1:46-53.

Kudithipudi D, Aguilar-Simon M, Babb J, Bazhenov M, Blackiston D, Bongard J, et al. Biological underpinnings for lifelong learning machines. *Nat Mach Intell.* 2022;4:196-210.

Ciprijanovic A, Kafkes D, Downey K, Jenkins S, Perdue GN, Madireddy S, et al. DeepMerge II. Building robust deep learning algorithms for merging galaxy identification across domains. *Mon Not R Astron Soc.* 2021;506:677-91.

Lao L, Kruger SE, Akcay C, Balaprakash P, Bechtel T, Howell E, et al. Application of machine learning and artificial intelligence to extend EFIT equilibrium reconstruction. *Plasma Phys Control Fusion.* 2022;64:074001.

Attarian S, Morgan D, Szlufarska I. Thermophysical properties of FLiBe using moment tensor potentials. *J Phys Chem B.* 2022;126:8618-28.

Jacobs R, Stetina K, Liu Y, Balaprakash P, Madireddy S, Morgan D. Deep learning of experimental electrochemistry for battery cathodes. *Appl AI Lett.* 2021;2:e28.

Madireddy S, Ramakrishna KS, Balaprakash P, Cerniglia D, Butler R, Warren R. Machine learning assisted development of a new probabilistic learning material discovery platform and its application in the design of high entropy alloys. *Mater Des.* 2021;205:109712.

Madireddy S, Balaprakash P, Karniadakis GE. Bayesian differential programming for robust systems identification under uncertainty. *Proc R Soc A.* 2020;476:20200290.

Madireddy S, Balaprakash P, Karniadakis GE. Enhancing gray box identification of nonlinear models through evolutionary optimization. *Eng Optim.* 2020;52:1655-72.

Liu Y, Madireddy S, Balaprakash P. Multipoint active learning for reduced-cost materials screening. *Comput Mater Sci.* 2021;194:110414.

Madireddy S, Balaprakash P, Balasubramanian G. Bayesian optimization of distributed neurodynamical models. *Mach Learn Sci Technol.* 2021;2:035006.