

ORIGINAL RESEARCH

Open access

Scientific Risk in Autonomous Materials Discovery: A Conceptual Framework for Anticipating Unsafe Recommendations

Hiroshi Nakamura^{1*}, Yuta Kato¹

Abstract

The integration of artificial intelligence into materials science has enabled autonomous discovery processes that accelerate the identification of novel compounds and structures. However, this advancement introduces scientific risks, including recommendations that may lead to unintended consequences, such as material instability, environmental hazards, or inefficiencies in application. This conceptual paper develops a framework for anticipating these risks by examining interaction dynamics between algorithmic outputs and systemic factors in research ecosystems. Drawing on recent literature, it synthesizes insights into how AI-driven autonomy influences epistemic reasoning and ethical trade-offs in materials discovery. The framework emphasizes steering logics that incorporate feedback structures for risk assessment, highlighting analytical implications for balancing innovation speed with precautionary measures. Through conceptual interpretations of uncertainty propagation and bias amplification, it explores how autonomous systems can inadvertently prioritize short-term optimization over long-term viability. Systems-level insights reveal the need for integrative approaches that align computational recommendations with broader societal and ecological considerations. Ultimately, this work underscores the importance of interpretive vigilance in AI-assisted discovery, offering a pathway to enhance resilience against unsafe outcomes while fostering sustainable progress in applied artificial intelligence for materials science.

Keywords Feedback structures, Uncertainty dynamics, Autonomous materials discovery, Scientific risk, AI recommendations, Ethical trade-offs

*Correspondence:

Hiroshi Nakamura

hiroshi.nakamura@gmail.com

¹ Department of Intelligent Materials Engineering, Faculty of Engineering, Nagoya University, Nagoya, Japan

Introduction

The advent of artificial intelligence (AI) in materials science is a transformative shift in how researchers discover and develop new materials. Traditional methods, reliant on iterative experimentation and human intuition, have given way to autonomous systems capable of processing vast datasets, simulating complex interactions, and generating recommendations at unprecedented scales [1, 2]. This evolution is particularly evident in autonomous materials discovery, where AI algorithms integrate with robotic platforms to conduct high-throughput screenings and predictive modeling, thereby compressing timelines from years to months [3]. Yet, as these systems gain autonomy, they introduce layers of scientific risk that extend beyond mere technical inaccuracies to encompass broader implications for research integrity and societal impact.

Scientific risk in this context refers to the potential for AI-generated recommendations to propagate uncertainties or biases that compromise the validity or safety of the materials discovered. For instance, autonomous discovery platforms may recommend compounds that appear optimal under simulated conditions but exhibit instability or toxicity when synthesized

in real-world environments [4]. Such risks arise from the interplay between algorithmic decision-making and the inherent complexities of material behaviors, where incomplete data representations can lead to misleading outcomes [5]. The conceptual exploration here focuses on anticipating these unsafe recommendations by understanding interaction dynamics, such as those between predictive models and external validation processes.

The motivation for this framework stems from the rapid proliferation of AI in materials science, driven by advancements in machine learning architectures and data availability [6]. Studies have demonstrated how AI can uncover novel alloys or nanomaterials with enhanced properties, such as improved conductivity or durability [7]. However, this acceleration comes with trade-offs: the opacity of deep learning models can obscure the rationale behind recommendations, fostering epistemic uncertainties that challenge traditional scientific scrutiny [8]. Moreover, ethical reasoning becomes paramount as autonomous systems influence resource allocation and environmental footprints in materials research [9].

At a systems level, the risks manifest through feedback loops in which initial algorithmic biases amplify across iterative cycles, potentially steering discovery toward suboptimal or hazardous paths [10]. For example, if training datasets overrepresent certain material classes, such as oxides, the system may undervalue alternatives, such as disordered alloys, limiting diversity in recommendations [11]. This not only affects scientific progress but also raises concerns about equitable innovation, where underrepresented materials could hold keys to sustainable applications [12].

Interpretive analysis reveals that unsafe recommendations often emerge from mismatches between computational ideals and physical realities. AI models trained on equilibrium assumptions may overlook kinetic barriers or environmental interactions, leading to suggestions that fail in practice [13]. Analytical implications extend to how researchers interpret these outputs: over-reliance on AI can erode human oversight, while under-utilization misses opportunities for efficiency [14]. Thus, the framework proposed herein seeks to integrate these dynamics, providing a conceptual lens for anticipating risks without resorting to empirical validation.

The literature underscores the urgency of addressing these issues. Recent syntheses highlight that AI autonomy in materials discovery has outpaced governance mechanisms, prompting calls for enhanced risk assessment [15]. Ethical considerations, such as the responsibility for AI-induced errors, further complicate the landscape [16]. By focusing on conceptual interpretations rather than predictive claims, this paper contributes to a nuanced understanding of how steering logics—balancing exploration and exploitation—can mitigate unsafe pathways.

In synthesizing theoretical backgrounds, the paper draws on interdisciplinary insights from AI ethics and materials informatics, revealing gaps in current approaches [17]. For instance, while uncertainty quantification techniques have advanced, their integration into autonomous workflows remains inconsistent [18]. The framework addresses this by emphasizing interaction dynamics that foster adaptive responses to emerging risks.

Ultimately, this conceptual work aims to guide researchers toward more resilient practices in autonomous materials discovery. By anticipating unsafe recommendations through systems-level insights and ethical reasoning, it promotes a balanced trajectory where AI enhances rather than undermines scientific integrity. The following sections delve into the theoretical foundations and articulate the proposed framework, setting the stage for deeper analytical implications in subsequent discussions.

Theoretical Background and Literature Synthesis

Evolution of autonomous systems in materials science

The integration of artificial intelligence into materials science has undergone a marked transition from assistive computational tools to self-directed discovery systems, fundamentally altering how scientific knowledge is generated and validated. Early applications of AI in the field primarily served as analytical accelerators, employing supervised and semi-supervised machine learning models to predict material properties from curated databases and simulation outputs [19].

These systems operated within a human-dominated epistemic loop, in which researchers retained primary authority over hypothesis formulation, interpretation, and experimental validation.

Advances in neural network architectures—particularly graph-based and message-passing models—enabled richer representations of atomic structures, bonding environments, and compositional spaces [20]. This representational shift expanded AI's role from pattern recognition to structural reasoning, enabling models to infer relationships previously inaccessible to linear descriptors. As a result, AI systems began to influence not only prediction but also search direction, shaping which regions of material space were explored and which were deprioritized.

This evolution culminated in the emergence of autonomous laboratories, where AI systems coordinate experimental design, synthesis, characterization, and iterative optimization with minimal human intervention [3]. In these environments, algorithms generate hypotheses, select experimental parameters, analyze outcomes, and update models in closed-loop cycles. Such systems dramatically compress discovery timelines and enable exploration at scales unattainable through conventional experimentation [1]. However, autonomy also introduces a qualitative shift in epistemic dynamics: decision authority increasingly migrates from human judgment to algorithmic steering mechanisms.

Interaction dynamics in autonomous materials systems are characterized by tight coupling between computational inference and physical execution. Real-time feedback between simulations, robotic platforms, and learning algorithms allows systems to adapt experimental strategies dynamically [21]. While this adaptability enhances efficiency, it also narrows interpretive space. Optimization-driven feedback can privilege rapid convergence toward locally optimal solutions, reducing opportunities for serendipitous discovery or for identifying anomalous behaviors that have traditionally informed scientific insight [22].

From a systems perspective, these developments amplify epistemic uncertainty in subtle ways. Classical materials discovery relied on human-mediated sense-making, where anomalies, inconsistencies, and failures were often epistemically productive. In contrast, autonomous systems prioritize consistency and performance metrics, rendering uncertainty increasingly data-centric and model-dependent [23]. When training datasets reflect historical biases—such as overrepresentation of stable or well-studied compounds—autonomous systems may reproduce and reinforce these preferences, shaping discovery trajectories in ways that appear rational yet remain epistemically narrow [11]. This shift sets the stage for new forms of scientific risk that cannot be reduced to model error alone.

Risks and uncertainties in AI-Driven discovery

Scientific risk in autonomous materials discovery arises from uncertainties intrinsic to AI-mediated inference, including representational approximations, dataset incompleteness, and extrapolative reasoning [24]. Unlike traditional experimental uncertainty, these risks propagate structurally through iterative decision loops. When predictive models—such as AI-augmented density functional theory surrogates—operate beyond their validated domains, uncertainty is not merely increased but transformed into actionable recommendations with apparent confidence [13].

A critical concern is uncertainty propagation, in which early-stage approximations influence downstream decisions, thereby compounding epistemic fragility across discovery cycles [25]. In autonomous workflows, each recommendation becomes new training data or validation evidence, allowing initial misrepresentations to cascade through successive iterations. This dynamic challenges conventional notions of error correction, as feedback loops may reinforce rather than attenuate epistemic distortions.

Bias amplification constitutes a second major risk pathway. Training on imbalanced or historically skewed datasets can systematically favor certain material classes, synthesis routes, or property regimes [26]. For example, databases dominated by oxide materials may bias discovery systems against chalcogenides or metastable phases, not because of inferior performance but because of representational scarcity [27]. Over time, such biases reshape exploration topology,

narrowing the range of materials considered viable. From an ethical standpoint, this raises concerns about inclusivity and sustainability, as overlooked materials may offer safer or more environmentally benign alternatives [12].

Feedback structures intrinsic to autonomous discovery further intensify these risks. Without explicit interpretive safeguards, systems may converge on solutions that appear robust within computational metrics yet conceal latent instabilities, toxicities, or synthesis challenges [28]. Recent work emphasizes the importance of hybrid epistemic architectures, where human judgment re-enters the loop not as a corrective afterthought but as an active interpretive agent capable of contextual reasoning [14]. Absent such integration, autonomous systems risk optimizing within epistemically impoverished spaces.

Ethical and epistemic dimensions of autonomy

Beyond technical uncertainty, autonomous materials discovery raises profound ethical and epistemic questions about responsibility, agency, and the legitimacy of knowledge. As AI systems increasingly shape experimental decisions, responsibility for unsafe or harmful outcomes becomes distributed across human designers, institutional infrastructures, and algorithmic processes [16]. This diffusion complicates accountability, particularly when recommendations result in environmentally hazardous materials or unsafe deployment pathways [9].

Ethical reasoning in this context cannot be relegated to post hoc evaluation. Conceptual analyses suggest that ethical steering must be embedded within discovery loops, influencing candidate selection, optimization objectives, and termination criteria [29]. Without such integration, autonomous systems may systematically favor materials that optimize narrow performance metrics while externalizing environmental or societal costs.

Epistemically, autonomy alters how knowledge is produced and validated. Over-automation risks diminishing researcher agency by substituting interpretive judgment with algorithmic authority [8]. Trade-offs emerge between speed and scrutiny: accelerated discovery may outpace verification processes, undermining confidence in results and eroding reproducibility [30]. Systems-level analyses, therefore, advocate for frameworks that integrate ethical evaluation, visibility into uncertainty, and human oversight as co-evolving elements of autonomous workflows [17].

While recent literature increasingly recognizes these challenges, significant gaps remain—particularly regarding long-term impacts on resource use, environmental sustainability, and global research equity [15]. These gaps underscore the need for conceptual models that treat ethics and epistemology as structural components of autonomous systems rather than external constraints.

Integration of risk anticipation strategies

Synthesizing these strands, the literature converges on the need for conceptual tools capable of anticipating unsafe recommendations before empirical failure occurs [31]. Risk anticipation in autonomous materials discovery depends on understanding interaction dynamics between algorithmic inference, uncertainty mediation, and human judgment. Adaptive feedback mechanisms—when explicitly designed—offer pathways for correcting bias, recalibrating uncertainty, and restoring epistemic balance [10].

Analytical implications point toward resilient discovery systems that balance autonomy with precautionary principles. Rather than constraining innovation, such systems enable responsible acceleration, ensuring that recommendations remain aligned with scientific validity, ethical accountability, and societal goals [5]. This synthesis motivates the conceptual framework advanced in the following section, which formalizes these dynamics into an integrated model of systemic risk formation and mitigation in autonomous materials discovery.

Proposed conceptual framework

The conceptual framework developed here provides an integrative, interpretive structure for anticipating the formation of scientific risk in autonomous materials discovery, with particular emphasis on how unsafe or misleading recommendations emerge from interactions within AI-driven systems. Rather than treating risk as a localized failure of prediction accuracy, the framework conceptualizes materials discovery as a coupled epistemic–computational system, in which algorithmic outputs, uncertainty representations, and normative constraints dynamically interact. Within this view, unsafe recommendations arise not from single errors but from system-level trade-offs, feedback amplification, and incomplete mediation of uncertainty and values.

At its core, the framework integrates three interdependent components:

- (1) algorithmic recommendation mechanisms,
- (2) uncertainty mediation layers, and
- (3) systemic oversight and steering loops.

Algorithmic recommendation mechanisms encompass the generative and predictive functions of AI systems, including surrogate models, generative architectures, and optimization pipelines that propose candidate materials based on encoded objectives. These mechanisms operate by transforming historical experimental data, simulations, or synthetic datasets into ranked material suggestions. However, their outputs are inherently shaped by objective formulations, representational biases, and limitations in the training data, rendering recommendations contingent rather than definitive.

Uncertainty mediation layers function as interpretive filters between raw algorithmic outputs and actionable scientific judgment. This layer captures how epistemic uncertainty—stemming from sparse data regimes, extrapolative inference, or model approximation—is quantified, visualized, or suppressed. Techniques such as Bayesian inference, ensemble modeling, or probabilistic calibration are used here. Yet, the framework highlights that uncertainty is often partially propagated or selectively attenuated, creating conditions in which apparently confident recommendations mask fragile epistemic foundations. In such cases, incomplete uncertainty mediation can lead to hazardous material recommendations, particularly when extrapolations extend beyond validated chemical or thermodynamic domains [24].

Systemic oversight and steering loops constitute the normative and corrective dimension of the framework. These loops integrate ethical reasoning, external validation protocols, and human judgment into the discovery process. Oversight does not function as a post-hoc constraint but as a continuous feedback structure that evaluates recommendations against safety, environmental, and societal criteria. Ethical considerations—such as toxicity, environmental persistence, or downstream misuse—are thus embedded as active interpretive constraints rather than external add-ons [9]. This component ensures that recommendation trajectories remain aligned with broader safety imperatives while preserving exploratory capacity.

Central to the framework are interaction dynamics that illustrate how recommendations evolve through iterative cycles. For example, an AI system may identify a high-performance alloy optimized for mechanical strength, yet without sufficient representation of environmental interactions or degradation pathways, such a recommendation may entail latent instability risks. Within the framework, this scenario is interpreted as a trade-off between short-term performance optimization and long-term material viability, revealing how narrow objective formulations can bias discovery pathways. The model therefore advocates steering logics grounded in multi-objective evaluation, where performance, stability, environmental compatibility, and safety co-evolve rather than compete in isolation [28].

Systems-level insights derived from the framework emphasize resilience through adaptive feedback structures. Iterative loops enable real-time recalibration when anomalies, biases, or epistemic inconsistencies are detected. For instance, evidence of bias amplification through recursive data reuse can trigger reweighting of objectives or expansion of uncertainty bounds, preventing premature convergence toward unsafe regions of the design space [10]. Ethical reasoning

is woven throughout these dynamics, promoting interpretations that account for societal consequences, such as minimizing reliance on toxic precursors or environmentally persistent compounds [5].

Collectively, the framework advances an interpretive shift in how safety is conceptualized in autonomous materials discovery. Risks are not treated as isolated technical failures but as emergent properties of interacting algorithmic, epistemic, and ethical processes. By explicitly mapping these interactions, the framework offers researchers a conceptual tool for anticipating unsafe trajectories before they materialize, enhancing scientific responsibility without constraining innovation. Systemic risk formation and its mediation through recommendation, uncertainty, ethics, and oversight pathways are illustrated in Figure 1.

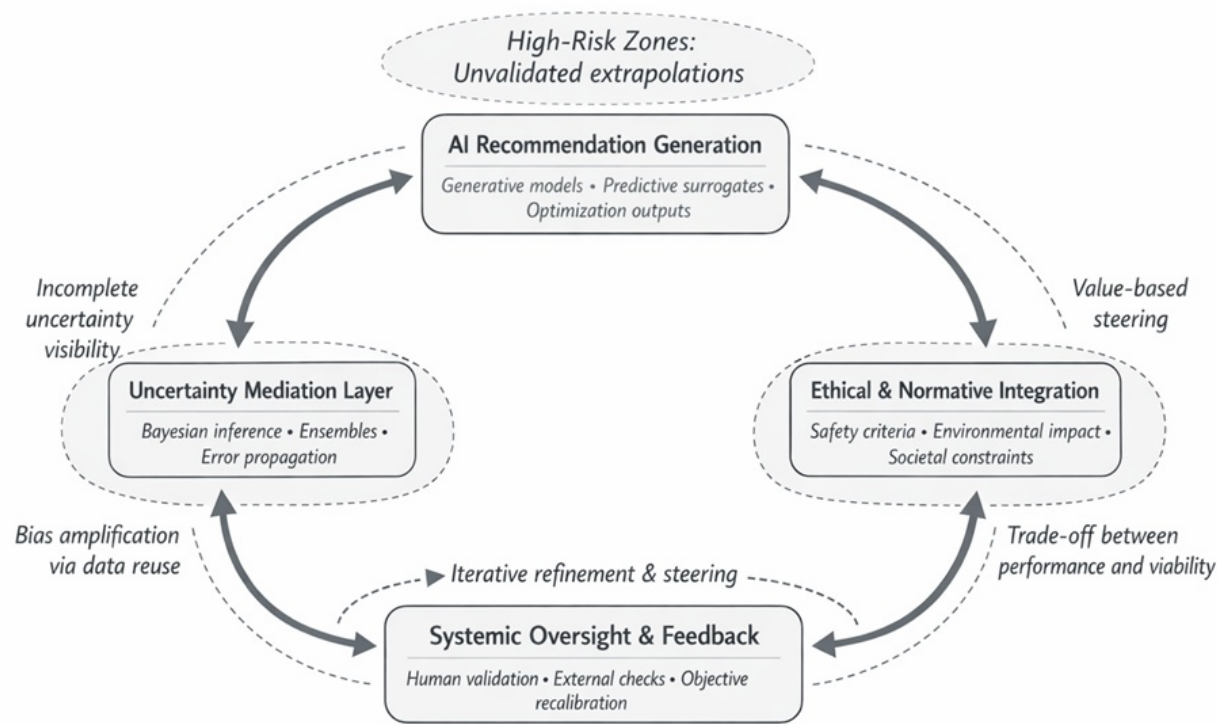


Figure 1. Systemic risk formation in autonomous materials discovery: an interpretive closed-loop oversight framework.

Analytical implications extend to how this framework informs practice, enabling researchers to anticipate unsafe paths by preemptively mapping dynamics. Conceptual interpretations underscore the need for integrative reasoning, in which risks are not isolated but emerge from system interactions [5]. By focusing on these elements, the framework offers a tool for enhancing safety without constraining innovation. To strengthen analytic traction, Table 1 translates the closed-loop oversight framework into a structured risk map that links (i) risk-generation mechanisms in autonomous discovery, (ii) observable manifestations in recommendation behavior, and (iii) corresponding steering and oversight interventions. This mapping clarifies how uncertainty propagation, bias stabilization, and ethical misalignment operate as coupled drivers of unsafe recommendations across autonomous workflows.

Table 1. Risk–mechanism–intervention map for anticipating unsafe recommendations in autonomous materials discovery

Framework locus (where risk forms)	Risk-generation mechanism (conceptual driver)	Typical manifestation in recommendations	Primary risk pathway (system dynamics)	Steering/oversight intervention (what prevents it)	Key refs

		(what “unsafe” looks like)			
Algorithmic recommendation mechanisms	Objective narrowing/performance monoculture: optimization targets encode short-term peaks and suppress long-horizon constraints	Candidate materials appear “optimal” in predicted metrics but fail under operating conditions (instability, incompatibility, impracticality)	Exploitation dominates exploration; model-led search converges to brittle optima	Explicit multi-objective formulations; constraint-aware optimization; diversification of acquisition policies	[5, 11, 28]
Algorithmic recommendation mechanisms	Representation blind spots: models encode equilibrium or simplified physics; kinetic/environmental variables omitted	High-confidence suggestions ignore synthesis feasibility, degradation pathways, or context sensitivity	Computational ideals decouple from physical reality; systematic omission of contingencies	Add feasibility gates; embed synthesis/processing constraints; enforce domain-of-validity checks	[5, 13, 20, 24]
Uncertainty mediation layer	Uncertainty attenuation: uncertainty exists, but is undercommunicated or collapsed into point estimates	Overconfident ranking; explanation narratives imply certainty; weak “risk visibility”	Uncertainty becomes non-actionable; decisions privilege velocity over epistemic humility	Calibrated UQ reporting; uncertainty visibility checkpoints; ensemble/Bayesian comparisons	[13, 18, 24]
Uncertainty mediation layer	Cascading error propagation: early-stage approximation shifts downstream search and validation priorities	Iterative cycles reinforce an initial misdirection; unsafe classes repeatedly resurfaced	Feedback amplification: an early error becomes a steering signal	Propagation-aware auditing; iterative recalibration triggers when variance grows; checkpointed retraining	[5, 24]
Systemic oversight and steering loops	Bias stabilization (beyond amplification): dataset imbalance becomes entrenched through iterative autonomy	Persistent over-recommendation of familiar classes (e.g., oxides) and under-exploration of alternatives	Feedback lock-in; narrowing of discovery topology over time	Bias-mitigating active learning; diversity constraints; entropy-targeted sampling; dataset balancing	[11, 22, 25, 26]
Systemic oversight and steering loops	Accountability diffusion: unclear responsibility for unsafe outcomes in autonomous pipelines	Safety-critical risks slip through because “the system recommended it”	Governance lags behind autonomy; validation becomes performative	Explicit responsibility allocation; mandatory human-in-the-loop gates for safety-critical decisions	[8, 9, 15, 28]
Ethical–epistemic integration	Normative misalignment: sustainability and hazard criteria treated as external, not embedded	Recommendations include toxic precursors or environmentally persistent options despite alternatives	Ethics as an afterthought; optimization ignores societal/ecological cost	Ethical filters inside the loop; disqualifying constraints; sustainability-aware steering	[9, 15, 17]

Scaling dynamics (system-level)	Nonlinear risk scaling with autonomy: coupling increases faster than oversight capacity	Unsafe convergence occurs faster at scale; issues emerge only after many cycles	Tight coupling + speed accelerates lock-in and oversights	Minimal working examples; staged autonomy; monitoring dashboards; gradual escalation of autonomy	[10, 14, 21, 28]
--	---	---	---	--	------------------

From error analysis to risk regimes

The principal analytical contribution of the proposed framework lies in reconceptualizing scientific risk in autonomous materials discovery as a system-level regime rather than a sequence of isolated prediction errors. Prior work has predominantly framed risk in terms of model accuracy, uncertainty magnitude, or data quality [5, 13, 24]. In contrast, the present framework demonstrates that unsafe recommendations emerge from interaction dynamics among algorithmic steering, uncertainty mediation, and oversight structures, positioning risk as an emergent property of coupled epistemic–computational systems [10, 15].

Risk as an emergent property of steering logics

A first analytical implication is that scientific risk is produced through steering logics embedded in autonomous discovery pipelines, rather than directly through model failure. Optimization routines, objective formulations, and reinforcement through iterative feedback cycles confer epistemic authority on algorithmic recommendations, even when underlying assumptions remain weakly validated [8, 28]. When steering mechanisms prioritize narrow performance targets—such as maximizing predicted strength or conductivity—risk accumulates structurally, favoring short-term optimization at the expense of long-term material stability or environmental compatibility [5, 11]. This reframing challenges evaluation paradigms that equate safety solely with predictive confidence, revealing that well-calibrated systems may nonetheless converge on unsafe discovery trajectories [13].

Epistemic compression and the loss of physical contingency

The framework further reveals a process of epistemic compression, wherein complex physical behaviors are reduced to tractable representations suitable for autonomous inference. While uncertainty quantification techniques have advanced considerably [18, 24], their operational integration into autonomous workflows often involves selective attenuation to preserve decision velocity [14]. As a result, physical contingencies such as kinetic barriers, degradation mechanisms, or synthesis constraints are systematically underrepresented. Analytical implications follow: unsafe recommendations arise not because uncertainty is absent, but because uncertainty is functionally marginalized within steering dynamics, producing outputs that are internally coherent yet externally fragile [5, 20].

Redistribution of scientific responsibility

A further implication concerns the redistribution of epistemic and ethical responsibility within autonomous discovery systems. As AI-generated recommendations increasingly guide experimental decisions, traditional loci of scientific judgment—hypothesis evaluation, validation sufficiency, and risk assessment—become distributed across algorithmic, human, and institutional layers [8, 15]. The framework clarifies that responsibility for unsafe outcomes cannot be confined to downstream validation stages alone. Instead, risk anticipation must be embedded in iterative recommendation loops, where acceptance criteria, visibility of uncertainty, and ethical constraints actively shape discovery trajectories [9, 16]. This insight reframes oversight as an intrinsic component of knowledge production rather than an external corrective mechanism.

Bias stabilization through iterative autonomy

From a systems perspective, the framework extends existing discussions of bias amplification [22, 26] by introducing the concept of bias stabilization. Under autonomous iteration, initial dataset imbalances—such as overrepresentation of

specific material classes—are not merely propagated but progressively entrenched through feedback-driven reinforcement [11, 25]. Once stabilized, these biases reshape the exploration topology of discovery systems, constraining diversity and increasing the likelihood of overlooking safer or more sustainable alternatives [12, 27]. Analytically, this positions bias not only as a fairness concern but as a long-term risk driver in materials innovation ecosystems.

Scaling risk with autonomy

The framework also provides insight into how scientific risk scales nonlinearly with system autonomy. As discovery platforms expand in speed, scope, and integration with experimental infrastructure, small epistemic misalignments can cascade across iterative cycles, producing disproportionate downstream consequences [3, 21]. Analytical implications indicate that increasing autonomy without a proportional increase in uncertainty mediation, ethical steering, and adaptive oversight structures systematically increases the probability of unsafe convergence [10, 28]. Risk, in this sense, scales with system coupling rather than with model complexity or dataset size alone [15].

Toward reflexive autonomous discovery

Collectively, these analytical implications support a transition toward reflexive autonomous discovery systems, in which AI platforms are designed not only to generate material candidates but to interrogate the conditions under which their recommendations remain valid, safe, and societally acceptable. Anticipating unsafe recommendations thus becomes an epistemic function embedded in system design, rather than a post hoc evaluative step [14, 17]. By enabling this shift, the framework contributes a conceptual foundation for resilient materials discovery, where innovation acceleration is coupled with sustained interpretive vigilance and ethical alignment rather than deferred precaution [5, 9, 15].

Results and Discussion

The conceptual framework for anticipating unsafe recommendations in autonomous materials discovery invites a deeper examination of its integrative potential within existing research ecosystems. Interaction dynamics, as conceptualized, reveal how AI autonomy can both accelerate progress and introduce vulnerabilities, necessitating a balanced interpretive approach [19]. For instance, while autonomous platforms excel in data processing, their reliance on historical datasets can perpetuate systemic biases, interpreted as feedback loops that constrain innovation diversity [26].

Ethical reasoning plays a pivotal role in navigating these dynamics, highlighting trade-offs between rapid discovery and responsible stewardship [16]. Systems-level insights suggest that without embedded safeguards, recommendations may favor materials with short-term benefits but long-term risks, such as environmental persistence issues [30]. This underscores the importance of steering logics that incorporate multi-stakeholder perspectives, fostering alignments that prioritize safety without stifling creativity [28].

Conceptual interpretations further explore how uncertainty mediation influences epistemic confidence. In autonomous workflows, the propagation of model approximations can lead to overconfident outputs, where interpretive vigilance becomes essential for discerning viable paths [3]. Analytical trade-offs here involve balancing algorithmic efficiency with comprehensive validations, ensuring that feedback structures adapt to emerging discrepancies [10].

Moreover, the framework's emphasis on systems-level resilience addresses gaps in current literature, where autonomy often outpaces risk governance [6]. By integrating ethical and epistemic elements, it offers a pathway to mitigate unsafe outcomes, such as through adaptive loops that recalibrate based on real-time insights [14]. This integrative approach not only enhances precautionary measures but also promotes sustainable advancements in materials science [4].

Challenges persist in implementing such dynamics, particularly in interdisciplinary contexts where varying epistemic standards may clash [7]. However, the framework's conceptual flexibility allows for tailored applications, interpreting risks

as opportunities for refinement rather than barriers [23]. Overall, this discussion reinforces the value of interpretive reasoning in transforming autonomous discovery into a more robust and ethically grounded endeavor [31].

Conclusion

In synthesizing the conceptual elements of scientific risk in autonomous materials discovery, this framework underscores the criticality of interaction dynamics and feedback structures for anticipating unsafe recommendations. Through systems-level insights and ethical reasoning, it illuminates pathways to balance innovation with precautionary vigilance, ensuring AI-driven processes contribute to sustainable progress. Interpretive approaches reveal that risks emerge from misalignments between computational ideals and physical realities, advocating for integrative steering logics that enhance resilience. Ultimately, this work provides a foundational lens for researchers to navigate the complexities of AI in materials science, fostering practices that prioritize safety amid accelerating discovery.

Acknowledgements

None

Conflict of interest

None

Financial support

None

Ethics statement

None

Received: 03 Mar 2025

Revised: 04 Apr 2025

Accepted: 26 Apr 2025

Published online: 18 July 2025

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

MacLeod BP, Parlange FGL, Morrissey TD, Häse F, Roch LM, Dettelbach KE, et al. Self-driving laboratory for accelerated discovery of thin-film materials. *Sci Adv.* 2020;6(20):eaaz8867.

Sha W, Guo Y, Yuan Q, Tang H, Li Y, Du Y, et al. Artificial intelligence to power the future of materials science and engineering. *Adv Intell Syst.* 2020;2(4):1900143.

Tom G, Schmid SP, Baird SG, Kusne AG, Yu L, Sparks TD, et al. Self-driving laboratories for chemistry and materials science. *Chem Rev.* 2024;124(17):9631-75.

Abolhasani M, Kumacheva E. The rise of self-driving labs in chemical and materials sciences. *Nat Synth.* 2023;2:483-92.

Honarmandi P, Arróyave R. Uncertainty quantification and propagation in computational materials science and simulation-assisted materials design. *Integr Mater Manuf Innov.* 2020;9:103-24.

Pilania G. Machine learning in materials science: From explainable predictions to autonomous design. *Comput Mater Sci.* 2021;193:110360.

Goswami L, Deka MK, Roy M. Artificial intelligence in material engineering: A review on applications of artificial intelligence in material engineering. *Adv Eng Mater.* 2023;25(11):2300104.

Green ML, Maruyama B, Schrier J. Autonomous (ai-driven) materials science. *Appl Phys Rev.* 2022;9(3):031401.

Hermann E, Hermann G, Tremblay JC. Ethical artificial intelligence in chemical research and development: A dual advantage for sustainability. *Sci Eng Ethics.* 2021;27(4):48.

Baird SG, Sparks TD. What is a minimal working example for a self-driving laboratory? *Matter.* 2022;5(9):2703-18.

Jablonka KM, Jothiappan GM, Wang S, Smit B, Aspuru-Guzik A. Bias free multi-objective active learning for materials design and discovery. *Nat Commun.* 2021;12(1):2312.

Goga AS. Integrating artificial intelligence in nanomaterials science: Pathways to revolutionary materials discovery and design. Ethics and risks. In: *International Conference on Innovative Research*; 2024.

Tran K, Neiswanger W, Yoon J, Zhang Z, Liang X, Xing E, et al. Methods for comparing uncertainty quantifications for material property predictions. *Mach Learn Sci Technol.* 2020;1(2):025006.

Abolhasani M, Brown KA. Role of ai in experimental materials science. *MRS Bull.* 2023;48(5):448-56.

DeCost BL, Hattrick-Simpers JR, Trautt Z, Kusne AG, Choudhary A, Kalidindi SR, et al. Scientific ai in materials science: A path to a sustainable and scalable paradigm. *Mach Learn Sci Technol.* 2020;1(3):033001.

Youssef A, Nichol AA, Martinez-Martin N, Rao A, Rath VK, Gross CP, et al. Ethical considerations in the design and conduct of clinical trials of artificial intelligence. *JAMA Netw Open.* 2024;7(9):e2432902.

Li J, Lim K, Yang H, Ren Z, Raghavan N, Chen W, et al. Ai applications through the whole life cycle of material discovery. *Matter.* 2020;3(2):393-432.

Korolev V, Nevolin I, Protsenko P. A universal similarity based approach for predictive uncertainty quantification in materials science. *Sci Rep.* 2022;12(1):15216.

Bukkapatnam STS. Autonomous materials discovery and manufacturing (amdm): A review and perspectives. *IISE Trans.* 2023;55(6):665-77.

Epps RW, Volk AA, Reyes KG, Abolhasani M. Accelerated ai development for autonomous materials synthesis in flow. *Chem Sci.* 2021;12(17):6025-36.

Delgado-Licon F, Abolhasani M. Research acceleration in self-driving labs: Technological roadmap toward accelerated materials and molecular discovery. *Adv Intell Syst.* 2023;5(5):2200331.

Kumagai M, Ando Y, Tanaka A, Tsuda K. Effects of data bias on machine-learning-based material discovery using experimental property data. *Sci Technol Adv Mater Methods.* 2022;2(1):341-9.

Das M, Perez TC, Shetty D, Hiremath P, Naik N. An overview on the role of artificial intelligence in modern advancements of material science. *ES Gen.* 2024;2:1183.

Varivoda D, Dong R, Omee SS, Hu J. Materials property prediction with uncertainty quantification: A benchmark study. *Appl Phys Rev.* 2023;10(2):021409.

Zhang H, Chen WW, Rondinelli JM. ET-AL: Entropy-targeted active learning for bias mitigation in materials data. *Appl Phys Rev.* 2023;10(2):021403.

Mavrogiorgos K, Kiourtis A, Mavrogiorgou A, Pitsios S, Kyriazis D, Varvarigou T, et al. Bias in machine learning: A literature review. *Appl Sci.* 2024;14(19):8860.

Soldatov MA, Butova VV, Pashkov D, Butakova MA. Self-driving laboratories for development of new functional materials and optimizing known reactions. *Nanomaterials.* 2021;11(3):619.

Seifrid M, Pollice R, Aguilar-Granda A, Gomes G, Aldeghi M, Hickman RJ, et al. Autonomous chemical experiments: Challenges and perspectives on establishing a self-driving lab. *Acc Chem Res.* 2022;55(17):2454-66.

Adetunla A, Akinlabi E, Jen TC. Harnessing the power of artificial intelligence in materials science: An overview. In: *Conference on Science*; 2024.

Blanco-Gonzalez A, Cabezon A, Seco-Gonzalez A, Calleja S, Casado-Vela J, Cuevas-Zuñiga B, et al. The role of ai in drug discovery: Challenges, opportunities, and strategies. *Pharmaceuticals (Basel).* 2023;16(6):891.

Pyzer-Knapp EO, Pitera JW, Staar PWJ, Takeda S, Laino T, Sanders DP, et al. Accelerating materials discovery using artificial intelligence, high performance computing and robotics. *NPJ Comput Mater.* 2022;8(1):84.