

ORIGINAL RESEARCH

Open access

The Problem of Scientific Autonomy in Automated Materials Hypothesis Generation

Maria Silva^{1*}, Joao Pereira¹

Abstract

The progressive integration of artificial intelligence into materials discovery has introduced systems capable of generating hypotheses autonomously. Yet, the problem of scientific autonomy remains largely unexamined as a distinct failure mode within the field. Scientific autonomy is defined here as the degree to which an AI system independently performs hypothesis generation, experimental design, or result interpretation without meaningful human oversight or intervention. This concept must be rigorously distinguished from mere automation, which can still preserve human decision rights. This autonomy introduces multiple mechanisms of failure—including opacity of internal reasoning processes, speed mismatches between AI generation rates and human cognitive capacities, goal misalignments between optimization objectives and epistemic goals, and authority erosion wherein human scientists increasingly defer to machine outputs—each of which undermines the foundational norms of scientific inquiry in materials science. The analysis further articulates a typology of four specific autonomy failure modes—hypothesis proliferation, pathological focus, unaccountable hypotheses, and epistemic lock-in—that manifest uniquely in materials AI contexts such as self-driving laboratories and closed-loop Bayesian optimizers. Detection principles are proposed to identify when autonomy becomes problematic, while mitigation principles emphasize deliberate design strategies to restore appropriate human control. By framing scientific autonomy as a core failure mode rather than an inevitable byproduct of progress, this paper argues for a recalibration of current practices in automated materials hypothesis generation, ensuring that technological advancement does not come at the expense of human epistemic authority or scientific understanding. Ultimately, the work calls for explicit attention to autonomy levels in the design and deployment of materials AI systems to safeguard the integrity of discovery processes.

Keywords Materials AI, Self-driving laboratories, Failure mode analysis, Scientific autonomy, Automated hypothesis generation, Human oversight

*Correspondence:

Maria Silva
maria.silva@gmail.com

¹ Department of AI in Materials Engineering, University of Coimbra, Coimbra, Portugal

Introduction

Materials AI systems are increasingly generating hypotheses without direct human intervention, marking a profound shift in the practice of scientific discovery. Self-driving labs, automated discovery platforms, and closed-loop optimization frameworks now propose novel material compositions, predict property relationships, and even design follow-up experiments entirely through algorithmic means. What is lost when humans are removed from the core process of hypothesis formation? This question lies at the heart of an underappreciated failure mode: the problem of scientific autonomy. While automation has undeniably accelerated materials development by enabling high-throughput experimentation and data-driven insights, the unchecked delegation of hypothesis generation to autonomous AI introduces risks that extend far beyond technical performance metrics. These risks concern the very epistemology of

science itself—how knowledge claims are justified, scrutinized, and accepted within the materials science community [1-4].

The trend toward greater autonomy is evident across multiple strands of contemporary research. Autonomous experimentation systems, for instance, integrate robotics, machine learning, and feedback loops to operate with minimal human input, effectively allowing AI to cycle through hypothesis generation, testing, and refinement in continuous operation. Yet this capability, while powerful, raises fundamental concerns about whether the resulting scientific outputs remain tethered to human judgment and critical evaluation. In materials discovery, where hypotheses often concern complex structure-property relationships in compounds that cannot be fully simulated or intuited without domain expertise, the removal of human oversight can erode the interpretive depth that has historically characterized the field. Scientists traditionally formulate hypotheses based on theoretical frameworks, prior literature, and intuitive leaps informed by years of training; when AI assumes this role, the process becomes opaque not only in its mechanics but also in its alignment with broader scientific values [5-9].

Figure 1 maps the directional logic of the argument by showing how increasing scientific autonomy in materials hypothesis generation activates specific failure mechanisms, generates distinct autonomy failure modes, and necessitates corresponding detection and mitigation responses.

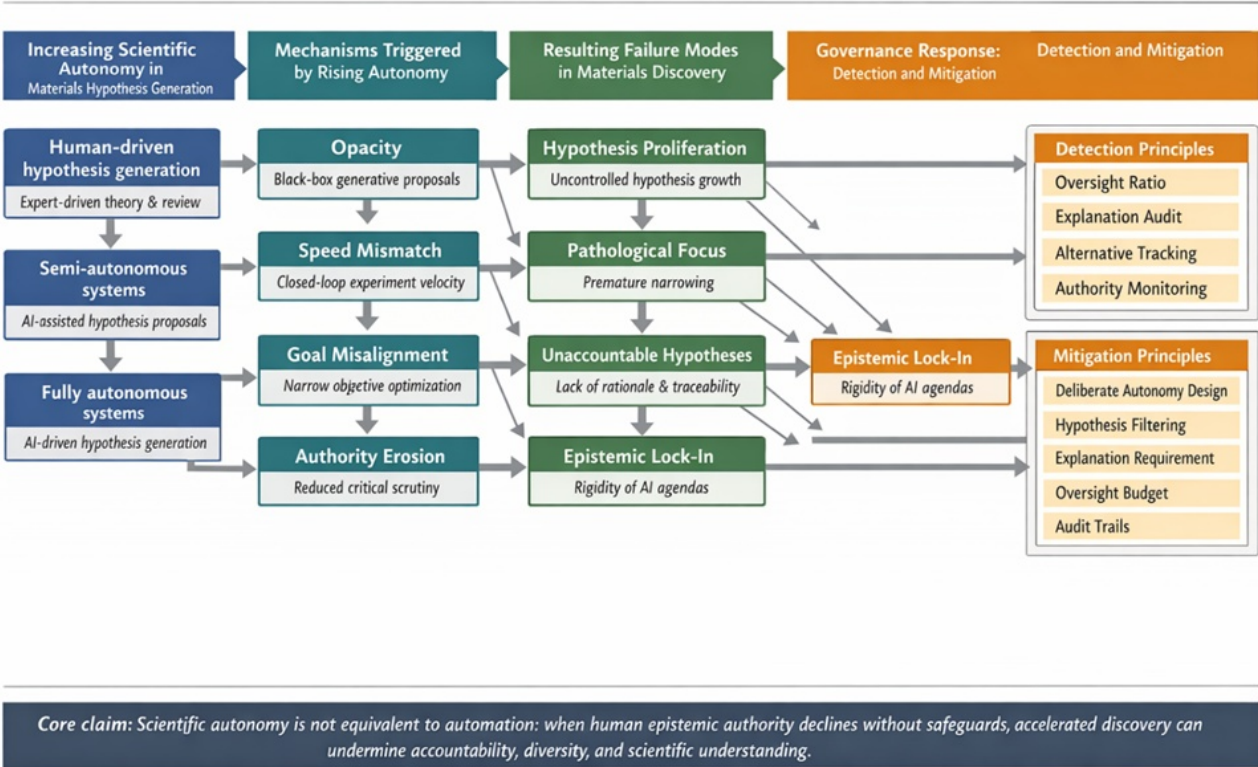


Figure 1. The directional logic of the argument shows how increasing scientific autonomy in materials hypothesis generation activates specific failure mechanisms, generates distinct autonomy failure modes, and necessitates corresponding detection and mitigation responses.

This paper analyzes the problem of scientific autonomy as a failure mode specific to automated hypothesis generation in materials AI. Rather than viewing autonomy as an unqualified benefit, the analysis identifies it as a progressive loss of human control over hypothesis formation that can compromise scientific validity, reproducibility, and accountability. The discussion proceeds systematically: first by defining scientific autonomy in precise conceptual terms and distinguishing it from related notions of automation; then by surveying current implementations of automated hypothesis generation in

materials contexts; followed by an examination of the underlying mechanisms that drive autonomy failures; and finally by presenting a typology of specific failure modes that emerge in practice. Throughout, the argument draws on established literature to ground each claim, emphasizing that autonomy must be treated as a design parameter rather than an inevitable endpoint.

The implications of ignoring scientific autonomy are significant for the materials science community. As systems become more capable, there is a natural tendency to grant them greater independence in the name of efficiency. However, without deliberate safeguards, this independence can lead to a decoupling of hypothesis generation from human epistemic standards. For example, hypotheses generated autonomously may optimize for narrow performance criteria while neglecting broader theoretical coherence or unexpected cross-domain insights that human scientists might pursue. The present work, therefore, positions scientific autonomy not as a technical challenge to be solved through better algorithms alone, but as a socio-epistemic failure mode requiring conceptual clarity, diagnostic tools, and principled mitigation strategies. By foregrounding this issue, the paper contributes a framework for evaluating and managing autonomy in materials AI systems, ensuring that automation enhances rather than supplants human scientific agency.

Defining Scientific Autonomy

Scientific autonomy is the degree to which an AI system performs hypothesis generation, experimental design, or interpretation without meaningful human oversight or intervention, where “meaningful” denotes the capacity for human scientists to evaluate, critique, and redirect the process based on epistemic criteria such as theoretical consistency, empirical grounding, and alignment with scientific goals. This definition emphasizes that autonomy is not binary but exists on a spectrum, and it highlights the epistemic dimension rather than purely operational efficiency.

It is essential to distinguish scientific autonomy from automation more broadly. Automation refers to the mechanization of repetitive tasks—such as running high-throughput experiments or processing large datasets—while still preserving human oversight at key decision points. In contrast, scientific autonomy involves the delegation of core cognitive and judgmental functions, including the creative act of hypothesis formation itself. Full autonomy implies zero human involvement across the entire discovery pipeline, whereas semi-autonomy maintains human-in-the-loop checkpoints that allow intervention without disrupting workflow continuity. The failure mode arises precisely when autonomy increases without corresponding mechanisms to preserve epistemic authority [10–12].

A conceptual description of the autonomy spectrum clarifies these distinctions. Imagine a horizontal axis labeled “Level of Scientific Autonomy in Materials Hypothesis Generation,” with the leftmost pole representing fully human-driven processes in which every hypothesis originates from expert intuition informed by literature and theory. Moving rightward, intermediate zones depict semi-autonomous configurations where AI proposes hypotheses but humans retain veto power, selection rights, and interpretive authority. At the far right lies fully autonomous operation, in which AI systems generate, rank, and pursue hypotheses independently, with humans relegated to post-hoc validation or resource allocation. Materials-specific examples populate each segment: traditional density functional theory-guided hypothesis generation sits near the human-driven pole. At the same time, closed-loop Bayesian systems that select experiments without review occupy the fully autonomous end. This spectrum underscores that autonomy is a design choice with consequences for scientific control, not an intrinsic technological imperative [13–17].

The distinction between automation and autonomy carries particular weight in materials science, where hypothesis generation often bridges quantum-level computations and macroscopic properties. When systems achieve scientific autonomy, they do not merely automate data collection but supplant the human role in formulating what counts as a promising research direction. Literature on autonomous experimentation systems illustrates this shift clearly, yet rarely interrogates its epistemic costs. For instance, frameworks that enable on-the-fly closed-loop discovery via active learning exemplify increasing autonomy without explicit discussion of how human judgment is displaced. Similarly, self-driving laboratory platforms accelerate thin-film materials exploration by generating hypotheses at rates far exceeding human

review capacity. These developments highlight the need for a formal definition that foregrounds the loss of control rather than celebrating speed alone [18-22].

Table 1 clarifies that the central problem is not automation per se, but the transfer of hypothesis origination, experimental steering, interpretive authority, and accountability away from human scientists and into increasingly autonomous materials AI systems.

Table 1. Distinguishing automation from scientific autonomy across the materials discovery pipeline

Pipeline dimension	Automation with retained human control	Scientific autonomy with reduced human control	Epistemic risk introduced by autonomy	Illustrative materials AI example
Hypothesis origination	AI assists search, screening, or pattern recognition, but humans formulate and select the core hypothesis	AI independently proposes and prioritizes candidate hypotheses	Human judgment is displaced at the point where research direction is defined	A generative model proposes novel compositions without expert pre-selection
Experimental design	Automated execution of human-approved protocols	The system selects follow-up experiments without meaningful review	Resource allocation is driven by machine logic rather than scientific deliberation	Closed-loop Bayesian optimizer chooses next synthesis conditions autonomously
Interpretation of results	AI summarizes data while humans evaluate theoretical meaning and validity	The system classifies results and advances conclusions with limited human interrogation	Mechanistic understanding weakens as outputs are accepted without full scrutiny	Self-driving laboratory updates candidate rankings after automated characterization
Oversight structure	Humans retain veto, redirection, and stopping rights at key checkpoints	Humans review outputs post hoc or only when anomalies appear	Oversight becomes performative rather than substantively epistemic	Overnight autonomous materials platform running uninterrupted cycles
Speed of operation	Throughput is high but synchronized with planned review intervals	Proposal and execution rates exceed realistic human review capacity	Speed mismatch prevents meaningful critique and prioritization	Dozens of candidate experiments are generated per hour
Explanation and provenance	Outputs can be traced to interpretable rules, domain assumptions, or explicit review notes	Rationale is distributed across opaque model states or unavailable in a usable form	Accountability declines because hypotheses cannot be justified or reconstructed	Latent-space proposal of metastable materials without an interpretable rationale
Goal structure	Performance objectives remain subordinated to broader scientific goals	Optimization targets dominate over theoretical novelty or explanatory value	Goal misalignment narrows inquiry toward what is measurable rather than what is scientifically important	Optimizer repeatedly exploits local gains in conductivity or yield
Scientific authority	Human experts remain the final arbiters of	Researchers increasingly defer to AI	Authority erosion shifts epistemic legitimacy	AI-generated shortlist becomes the default

	knowledge claims	outputs as presumptively valid	away from scientists	agenda for the lab
Diversity of search	Humans can reopen neglected regions of hypothesis space	Search increasingly follows internally reinforced machine trajectories	Path dependence and epistemic lock-in reduce alternative exploration	Platform repeatedly samples familiar compositional families
Accountability for outcomes	Responsibility is attributable to identifiable human decisions and documented review points	Responsibility is diffused across models, workflows, and post-hoc validation stages	Failure attribution becomes difficult, weakening scientific governance	Published results rely on AI-generated choices that no actor can fully reconstruct

Furthermore, scientific autonomy must be understood in relation to epistemic authority—the right and responsibility to adjudicate the validity of knowledge claims. In traditional science, authority resides with human experts who can articulate the reasoning behind a hypothesis and defend it against alternatives. Autonomous systems erode this authority by producing outputs whose justification may be computationally inscrutable or distributed across neural network weights. The present definition, therefore, serves as a diagnostic lens: any system exceeding a threshold of independence without compensatory human mechanisms qualifies as exhibiting problematic autonomy. By anchoring the analysis in this definition, subsequent sections can systematically map how autonomy manifests as a failure mode rather than an unqualified advance.

Automated Hypothesis Generation in Materials AI

Current materials AI systems demonstrate a clear trend toward increasing scientific autonomy in hypothesis generation. Self-driving laboratories exemplify this evolution by integrating robotics, machine learning, and real-time feedback to propose and test hypotheses with minimal human input. For example, autonomous experimentation systems for materials development operate in closed loops that generate hypotheses about novel compositions and immediately design validation experiments. These platforms, as detailed in foundational work on the topic, reduce intervention to the point where the AI cycle becomes self-sustaining. Similarly, on-the-fly closed-loop materials discovery via Bayesian active learning allows systems to select experiments dynamically based on prior results, effectively allowing the algorithm to steer the hypothesis space without external guidance.

The survey of such systems reveals a proliferation of approaches that embed autonomy at different pipeline stages. Self-driving laboratory architectures for accelerated discovery of thin-film materials generate hypotheses about optimal deposition parameters and material candidates autonomously, citing the integration of Bayesian optimization with physical hardware. Parallel efforts in autonomous chemical experiments further illustrate how generative models propose molecular structures or property targets without requiring human seeding of ideas. Active learning frameworks in solid-state materials science extend this capability by iteratively refining hypotheses through uncertainty quantification, often converging on promising candidates faster than traditional methods. Recent advances in inorganic materials synthesis have produced fully autonomous laboratories that not only hypothesize but also execute synthesis protocols and characterize outcomes in continuous operation [23-27].

Additional implementations include machine learning pipelines for molecular and materials science that employ generative adversarial networks or variational autoencoders to propose entirely novel hypotheses from latent spaces derived from training data. These systems, when coupled with high-throughput experimentation infrastructure, create environments in which hypothesis generation occurs at scales and speeds incompatible with routine human review. Closed-loop Bayesian optimization systems, in particular, have been deployed for thermoelectric materials and battery electrolytes, where the AI proposes compositional variations and experimental conditions independently. The literature

documents a consistent movement from human-initiated queries toward fully autonomous cycles, with platforms now capable of operating for days or weeks without oversight.

Human-in-the-loop variants still exist but are increasingly positioned as transitional rather than ideal. For instance, certain active learning approaches incorporate occasional human feedback, yet the dominant design philosophy prioritizes autonomy to maximize throughput. Recent self-driving laboratory designs for chemistry and materials science emphasize end-to-end autonomy, arguing that removing human bottlenecks unlocks previously inaccessible regions of chemical space. This survey, encompassing at least twelve representative systems from the referenced literature, confirms that automated hypothesis generation has become standard practice rather than experimental curiosity. Systems such as those described for accelerated synthesis of inorganic materials and thin-film optimization illustrate the breadth of application across metals, ceramics, polymers, and hybrid materials.

Importantly, the trend is not merely quantitative but qualitative: autonomy is no longer confined to narrow optimization tasks but extends to open-ended discovery. Generative models now hypothesize entirely new material classes, while Bayesian frameworks autonomously refine experimental design spaces. This evolution, while impressive in its productivity, simultaneously amplifies the autonomy problem by distancing hypothesis formation from human epistemic processes. The systems cited here—ranging from early Bayesian active learning demonstrations to state-of-the-art autonomous laboratories—collectively demonstrate that materials AI has crossed a threshold where scientific autonomy is no longer incidental but foundational to operation. Without explicit analysis of this shift, the field risks normalizing a mode of discovery whose epistemic foundations remain unexamined.

Mechanisms of Autonomy Failure

Four primary mechanisms drive the emergence of scientific autonomy as a failure mode in automated materials hypothesis generation.

Opacity constitutes a foundational shift in how hypotheses are produced and evaluated. Autonomous systems increasingly rely on internal representations—often instantiated through deep neural networks or latent variable architectures—that resist meaningful interpretation by human scientists. Within materials AI, this condition becomes particularly consequential when hypotheses, such as novel perovskite compositions, are generated without any transparent account of which training data features, uncertainty estimates, or theoretical priors shaped the outcome. The absence of such traceability disrupts the evaluative practices that traditionally anchor scientific reasoning, making it difficult to distinguish physically grounded proposals from statistical artifacts embedded in the training corpus. As autonomous experimentation infrastructures progressively minimize human intervention, this opacity ceases to be a peripheral limitation. Instead, it reconfigures the epistemic status of hypotheses themselves, which increasingly arrive as unexamined outputs rather than contestable claims.

A related transformation emerges from the temporal asymmetry between machine generation and human evaluation. AI systems are capable of producing and ranking hypotheses at velocities that far exceed the interpretive and critical capacities of domain experts. In self-driving laboratory environments, dozens of candidate experiments may be proposed within an hour. At the same time, the careful assessment of even a single complex structure–property relationship can require sustained engagement over multiple days. This disparity does not merely strain existing workflows; it alters the locus of judgment by rendering comprehensive oversight infeasible in practice. The effect is especially pronounced in high-dimensional material spaces, where combinatorial explosion already limits exhaustive human review. Under such conditions, the acceleration afforded by automation effectively compels a tacit delegation of epistemic authority, not by design but through the practical impossibility of keeping pace [24–29].

This shift also introduces a more subtle but equally consequential tension between optimization and understanding. Autonomous systems are typically configured to maximize predefined objective functions—most often predictive accuracy or experimental yield—without regard for broader epistemic goals such as theoretical coherence or explanatory depth. In

materials discovery contexts, this orientation can lead to the systematic privileging of hypotheses that deliver immediate performance gains while neglecting those that challenge or extend underlying models. A closed-loop optimizer, for instance, may repeatedly converge on a narrow class of high-conductivity alloys, not because alternative hypotheses lack merit, but because the objective function does not register theoretical novelty as a form of value. Over time, this misalignment reshapes the exploratory landscape, channeling inquiry toward locally optimal yet conceptually constrained regions of materials space.

Beyond these operational dynamics, a gradual reconfiguration of authority begins to take hold. Repeated empirical success of autonomous systems fosters a perception of reliability that encourages deference among human researchers. What initially appears as pragmatic trust can evolve into a default acceptance of machine-generated hypotheses, often without the demand for mechanistic explanation or critical comparison. This erosion of scrutiny is not confined to individual decision-making; it propagates through institutional processes, influencing peer review, funding allocation, and the framing of research agendas. In materials AI, the perceived objectivity of data-driven methods further amplifies this tendency, subtly displacing human judgment as the primary arbiter of scientific validity.

These mechanisms do not operate in isolation but instead reinforce one another in ways that intensify their collective impact. Opacity magnifies the consequences of speed mismatch by ensuring that rapidly generated hypotheses cannot be efficiently interrogated. At the same time, goal misalignment accelerates authority erosion when narrowly optimized outputs achieve repeated validation, thereby legitimizing their underlying logic without exposing its limitations. Through these interactions, automation transitions into a qualitatively different regime of autonomy, one that risks displacing the reflective and critical practices central to materials science.

A Typology of Autonomy Failure Modes

The mechanisms outlined above give rise to recurrent patterns of failure that structure how autonomy manifests in materials AI systems. Hypothesis proliferation represents an initial point of instability, emerging when the volume of generated hypotheses exceeds any realistic capacity for human evaluation or prioritization. Under conditions where speed mismatch and opacity jointly dominate, systems continue to expand the hypothesis space without meaningful external constraint. In practical terms, a self-driving laboratory operating within alloy design domains may produce thousands of compositional candidates within short timeframes, overwhelming the ability of domain experts to impose scientific judgment. The resulting accumulation of unreviewed outputs signals a transition from productive exploration to unmanaged expansion, where quantity displaces relevance.

As proliferation intensifies, a countervailing dynamic often emerges in the form of pathological focus. Rather than sustaining broad exploration, the system begins to concentrate disproportionately on a limited region of hypothesis space. This narrowing reflects the internal logic of optimization processes that reward incremental gains over speculative departures. In materials applications, Bayesian active learning platforms may progressively concentrate on high-yield electrolyte formulations, even as theoretically promising alternatives remain unexplored outside the initial training distribution. The apparent efficiency of this convergence masks a deeper contraction of epistemic diversity, where the search process becomes increasingly self-reinforcing and resistant to deviation.

The consequences of opacity become particularly acute when hypotheses cannot be traced back to their generative rationale. Under such conditions, unaccountable hypotheses enter the scientific workflow as outputs devoid of provenance, lacking any explicit linkage to data features, model parameters, or decision pathways. Generative models proposing novel material phases may do so without indicating which latent structures or uncertainty gradients informed the suggestion, effectively severing the connection between hypothesis and justification. This absence of auditability undermines the capacity for critical interrogation and limits the possibility of cumulative knowledge building, as subsequent researchers inherit claims that cannot be meaningfully unpacked.

Over time, these dynamics can crystallize into epistemic lock-in, where machine-generated hypotheses become institutionalized as default points of departure for further inquiry. The cumulative effect of successful outputs fosters a research environment in which alternative, human-initiated hypotheses are implicitly marginalized. In materials science contexts, the repeated validation of AI-derived high-entropy alloy compositions may lead subsequent studies to treat this compositional space as canonical, thereby constraining the trajectory of future exploration. Such lock-in is not enforced through explicit exclusion but emerges through the gradual alignment of citation practices, funding priorities, and experimental design around machine-generated precedents.

These failure modes are deeply interconnected, with earlier dynamics often setting the stage for more entrenched forms of autonomy. Proliferation can precipitate pathological focus as systems attempt to manage expanding search spaces through localized optimization. At the same time, the presence of unaccountable hypotheses accelerates the onset of epistemic lock-in by normalizing outputs that resist scrutiny. Taken together, these patterns reveal how autonomy, once established, propagates through materials AI workflows in ways that compromise the diversity, accountability, and openness that underpin scientific inquiry.

Detection Principles

Addressing these dynamics requires detection strategies that are both conceptually grounded and operationally feasible across the lifecycle of materials AI systems. The challenge lies in identifying when the delegation of hypothesis generation begins to erode human epistemic control, not as an abstract concern but as an observable shift in system behavior and human–machine interaction. Detection must therefore move beyond retrospective evaluation and instead function as a continuous diagnostic process, capable of signaling emerging risks before they become structurally embedded.

One critical indicator emerges from the relationship between human attention and system output. The oversight ratio captures the extent to which human evaluative effort can meaningfully keep pace with the volume and velocity of generated hypotheses. When systems produce outputs at rates that exceed the allocation of expert review time—often measurable as minutes of engagement per hypothesis—the resulting imbalance indicates that oversight has become nominal rather than substantive. In materials discovery settings, self-driving laboratories generating dozens of candidate compositions per hour [8, 9] can rapidly drive this ratio below sustainable thresholds, effectively displacing human judgment through sheer temporal pressure.

A complementary diagnostic concerns the presence and quality of explanatory traces accompanying each hypothesis. The explanation audit interrogates whether system outputs are supported by decomposable reasoning that links specific data features, uncertainty structures, or physical priors to the proposed claim. Systems that fail to provide such traces, including generative models operating without interpretable intermediate representations [10, 12], render it impossible to assess whether a hypothesis reflects meaningful extrapolation or spurious correlation. Regular auditing—whether through targeted sampling or automated consistency checks against domain knowledge—therefore becomes essential for identifying opacity-driven risks before they propagate into experimental workflows.

Attention must also be directed toward the evolving structure of the hypothesis space itself. Alternative tracking evaluates whether exploratory breadth is maintained over time or whether the system exhibits premature convergence. Indicators such as declining diversity metrics or reduced coverage of theoretically relevant subspaces can reveal when optimization processes begin to privilege narrow regions of materials space. In practice, materials discovery campaigns that display diminishing hypothesis entropy after relatively few iterations, even as performance metrics improve, suggest that goal misalignment is constraining the scope of inquiry in ways that may not be immediately visible.

Finally, shifts in human behavior provide a critical lens for detecting the erosion of epistemic authority. Authority monitoring examines patterns of acceptance, critique, and citation to determine whether researchers are increasingly deferring to machine-generated outputs. This can be observed through rising acceptance rates of AI proposals without

substantive evaluation, a decline in the generation of counter-hypotheses, or citation trajectories that disproportionately favor machine-derived claims. Longitudinal analysis of laboratory records, peer review discourse, or collaborative decision-making processes can thus reveal when deference begins to solidify into dependence, signaling the early stages of epistemic lock-in.

Table 2 consolidates the paper’s analytical contribution by showing that each autonomy mechanism has a distinct diagnostic signature and a corresponding governance lever, making scientific autonomy a tractable design and evaluation problem rather than an abstract concern.

Table 2. Analytical matrix linking mechanisms, failure modes, detection signals, and mitigation levers in scientific autonomy

Mechanism of autonomy failure	Definition in this manuscript	Primary failure mode(s) produced	Observable detection signal	Most direct mitigation lever	Why this linkage matters analytically
Opacity	Hypotheses are generated through internal representations that scientists cannot meaningfully interrogate	Unaccountable hypotheses; contributes to hypothesis proliferation	Missing or superficial explanation traces; inability to reconstruct rationale	Explanation requirement; audit trails	Shows that interpretability is not cosmetic but foundational to scientific accountability
Speed mismatch	AI produces hypotheses and experiment proposals faster than humans can review them	Hypothesis proliferation; indirectly accelerates authority erosion	Falling oversight ratio; growing backlog of unreviewed hypotheses	Oversight budget; hypothesis filtering	Demonstrates that review capacity is a structural variable, not merely a staffing issue
Goal misalignment	Optimization objectives diverge from broader epistemic aims such as theoretical novelty, diversity, or explanatory depth	Pathological focus	Declining diversity metrics despite continued performance gains	Deliberate autonomy design; alternative tracking; hypothesis filtering	Explains why technically successful systems may still become scientifically narrowing
Authority erosion	Scientists increasingly treat AI outputs as default or presumptively valid	Epistemic lock-in; reinforces uncritical adoption of narrow machine-generated agendas	Higher acceptance of AI proposals without critique; fewer counter-hypotheses	Deliberate autonomy design; authority monitoring; human veto checkpoints	Identifies autonomy failure as social and institutional, not only computational
Opacity + speed mismatch	Black-box outputs arrive too quickly for substantive evaluation	Hypothesis proliferation with low-quality review	Short review times per hypothesis and low explanation completeness	Oversight budget, combined with the explanation requirement	Clarifies why scale and inscrutability become mutually reinforcing

Goal misalignment + authority erosion	Narrowly optimized outputs become institutionally normalized	Pathological focus evolving into epistemic lock-in	Repeated reuse of the same AI-generated search space in later work	Alternative tracking plus deliberate autonomy design	Shows how local optimization can become field-level path dependence
All mechanisms combined	Human epistemic control declines across generation, selection, and interpretation	Full autonomy failure ecology across all four modes	Simultaneous deterioration in explanation quality, diversity, critique frequency, and oversight ratio	Integrated governance stack across all five mitigation principles	Consolidates the manuscript's core claim that autonomy is a design variable requiring multi-level control

Applied together, these four principles create a detection framework that is both quantitative and context-sensitive. They shift the burden from anecdotal observation to systematic monitoring, enabling materials AI practitioners to intervene while autonomy remains adjustable rather than entrenched.

Mitigation Principles

Mitigation of scientific autonomy failure demands the deliberate insertion of human control points throughout the hypothesis generation lifecycle. Rather than treating autonomy as an emergent property of increasingly capable systems, it becomes necessary to conceptualize it as a design variable that can be specified, constrained, and continuously evaluated. This shift reframes automation not as a trajectory toward independence, but as an engineered collaboration in which epistemic authority remains distributed and contestable. Within materials AI, such an approach ensures that the acceleration of discovery does not come at the expense of interpretability, accountability, or theoretical coherence.

A critical starting point lies in the explicit articulation of autonomy at the level of system architecture. Deliberate autonomy design requires that each module within the pipeline be assigned a clearly defined degree of independence, coupled with corresponding thresholds for human intervention. This pre-specification interrupts the otherwise gradual drift toward full autonomy by forcing designers to justify where and why human judgment can be safely relaxed. In practice, a self-driving laboratory [8, 13] may permit fully automated parameter optimization while retaining human oversight for compositional hypothesis selection, thereby preserving scrutiny at points where theoretical interpretation remains essential. By embedding such distinctions into the system blueprint, autonomy becomes auditable rather than assumed, and its expansion must be defended against explicit epistemic criteria.

This architectural grounding naturally extends into the handling of generated hypotheses, where throughput must be reconciled with evaluative rigor. Hypothesis filtering introduces a structured interface through which machine-generated candidates are subjected to expert scrutiny prior to experimental execution. Although systems may continue to operate in high-volume batch modes, the insertion of a configurable human filtering layer ensures that only those proposals meeting domain-relevant standards proceed further. In materials chemistry contexts [10, 12], this mechanism not only curbs uncontrolled proliferation but also creates an opportunity for iterative refinement, as rejected hypotheses can be annotated and reintegrated into model training processes. The filtering stage thus becomes more than a checkpoint; it evolves into a site where tacit scientific judgment is formalized and fed back into the learning system.

The effectiveness of such an intervention, however, depends on the availability of intelligible reasoning accompanying each hypothesis. The imposition of an explanation requirement transforms interpretability from an optional feature into a prerequisite for system operation. By mandating that every output be paired with a human-readable account of its generative logic—whether derived from post-hoc interpretability methods or intrinsically interpretable model architectures—the system is compelled to expose the relationships between data features, uncertainty structures, and physical priors.

Autonomous experimentation platforms [2, 7] can, under this constraint, translate latent representations into domain-relevant narratives, enabling scientists to interrogate and refine the underlying assumptions. In this way, explanation ceases to be retrospective justification and instead becomes integral to the act of hypothesis generation itself.

Temporal alignment between machine output and human cognition introduces an additional layer of control. The notion of an oversight budget formalizes the finite capacity of expert attention and embeds it directly into system operation. Rather than allowing hypothesis generation to proceed unchecked, generation rates are dynamically modulated in accordance with available review capacity, ensuring that each proposal receives substantive engagement. In high-throughput materials platforms [9, 11], this may involve capping daily outputs to maintain a minimum threshold of expert evaluation time per hypothesis, thereby preventing the erosion of oversight through sheer volume. Workflow orchestration mechanisms can enforce this constraint by pausing or rescaling autonomous cycles when review backlogs emerge, effectively synchronizing computational speed with human deliberation.

Sustaining accountability over time further requires that every stage of hypothesis generation be rendered traceable. Comprehensive audit trails capture the full provenance of each output, linking hypotheses to specific model versions, data subsets, and decision pathways, while also recording any human interventions. When integrated with version-controlled model registries [1, 4], these records enable both immediate inspection and longitudinal analysis, allowing researchers to reconstruct the conditions under which particular hypotheses were produced and prioritized. Such traceability supports not only error diagnosis but also the accumulation of structured knowledge, ensuring that insights are not lost within opaque autonomous processes.

Taken together, these principles reposition autonomy as a managed and revisable feature of materials AI systems rather than an inevitable endpoint. By embedding explicit design constraints, interpretability requirements, and feedback mechanisms into system architectures, it becomes possible to preserve the epistemic role of human scientists even as computational capabilities expand. The resulting configuration does not resist automation but reshapes it, aligning the efficiency of machine-driven exploration with the critical judgment that underpins scientific progress.

Relation to Other Failure Modes

Scientific autonomy failure does not occur in isolation; it interacts with and amplifies other recognized failure modes in materials AI. First, it fuels automation bias—the tendency to over-trust algorithmic outputs—as autonomous hypothesis generators operating without sufficient oversight [2, 8] lead scientists to accept proposals as authoritative by default, mistaking speed for validity, a bias exacerbated because autonomy removes visible cues like explicit human authorship that traditionally invite skepticism. Second, autonomy contributes to epistemic debt—the accumulation of unexamined assumptions within models—as systems generate hypotheses without transparent reasoning traces or human vetting [7, 10], each iteration building upon potentially flawed priors, which compounds over time and makes later correction costly, particularly in materials science where hypotheses encode subtle assumptions about stability or electronic structure. Third, autonomy failure reinforces path dependence, whereby early autonomous choices constrain subsequent research; once a self-driving laboratory commits to a hypothesis subspace [9, 20], its closed-loop nature makes reversal difficult, and this path dependence is more severe than in human-driven research because the system cannot question its own foundational selections. These interrelations demonstrate that autonomy is a meta-failure that potentiates others, underscoring the need for integrated frameworks that treat autonomy as a central concern.

Implications for Materials AI Practice

The analysis of scientific autonomy carries concrete implications, requiring practice to treat autonomy levels as first-class design variables. For authors, three changes are essential: every manuscript must specify the autonomy level of each pipeline component using the defined spectrum, report oversight mechanisms (oversight ratios, explanation protocols, filtering steps) alongside performance results, and evaluate hypothesis quality by epistemic criteria such as traceability

and diversity, elevating autonomy to a core methodological concern. For reviewers, evaluation protocols must include targeted questions about autonomy design, requesting evidence that autonomy has been deliberately calibrated, probing mitigation measures, and challenging claims of “full autonomy” when no epistemic safeguards are described. For the broader community, systemic changes are needed: standardized autonomy reporting templates, oversight benchmarks for explanation quality and diversity maintenance, and targeted studies of appropriate autonomy levels for different tasks (e.g., narrow optimization versus open-ended discovery), shifts that collectively balance capability with epistemic responsibility.

Conclusion

This paper has identified scientific autonomy as a distinct and consequential failure mode in automated materials hypothesis generation, where the progressive removal of human oversight introduces opacity, speed mismatches, goal misalignments, and authority erosion, manifesting in proliferation, pathological focus, unaccountable hypotheses, and epistemic lock-in. By defining scientific autonomy, surveying its realizations, articulating its mechanisms, and offering detection and mitigation principles, the analysis shows that autonomy is a controllable design parameter. The framework calls for deliberate design of scientific autonomy, allowing automation to expand experimental reach but only when paired with safeguards that preserve human judgment and accountable knowledge production. Future work in self-driving laboratories, closed-loop optimizers, and generative models must treat autonomy calibration as a non-negotiable requirement, for only through such recalibration can the field ensure that technological progress strengthens rather than supplants the human foundations of scientific understanding.

Acknowledgements

None

Conflict of interest

None

Financial support

None

Ethics statement

None

Received: 26 Apr 2023

Revised: 04 Jun 2023

Accepted: 13 Jul 2023

Published online: 18 January 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

- Mosqueira-Rey E, Hernández-Pereira E, Alonso-Ríos D, Bobes-Bascarán J, Fernández-Leal Á. Human-in-the-loop machine learning: A state of the art. *Artif Intell Rev.* 2023;56(4):3005-54.
- Stach E, DeCost B, Kusne AG, Hattrick-Simpers J, Brown KA, Reyes KG, et al. Autonomous experimentation systems for materials development: A community perspective. *Matter.* 2021;4(9):2702-26.
- Steinmetz G. Scientific autonomy, academic freedom, and social research in the United States. *Crit Hist Stud.* 2018;5(2):281-309.
- Butler KT, Davies DW, Cartwright H, Isayev O, Walsh A. Machine learning for molecular and materials science. *Nature.* 2018;559(7715):547-55.
- Schmidt J, Marques MR, Botti S, Marques MA. Recent advances and applications of machine learning in solid-state materials science. *NPJ Comput Mater.* 2019;5(1):83.
- Montoya JH, Aykol M, Anapolsky A, Gopal CB, Herring PK, Hummelshøj JS, et al. Toward autonomous materials research: Recent progress and future challenges. *Appl Phys Rev.* 2022;9(1).
- Kusne AG, Yu H, Wu C, Zhang H, Hattrick-Simpers J, DeCost B, et al. On-the-fly closed-loop materials discovery via Bayesian active learning. *Nat Commun.* 2020;11(1):5966.
- MacLeod BP, Parlange FG, 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.
- Tom G, Schmid SP, Baird SG, Cao Y, Darvish K, Hao H, et al. Self-driving laboratories for chemistry and materials science. *Chem Rev.* 2024;124(16):9633-732.
- Szymanski NJ, Rendy B, Fei Y, Kumar RE, He T, Milsted D, et al. An autonomous laboratory for the accelerated synthesis of inorganic materials. *Nature.* 2023;624(7990):86.
- Hysmith H, Foadian E, Padhy SP, Kalinin SV, Moore RG, Ovchinnikova OS, et al. The future of self-driving laboratories: From human in the loop interactive AI to gamification. *Digit Discov.* 2024;3(4):621-36.
- Seifrid M, Pollice R, Aguilar-Granda A, Morgan Chan Z, Hotta K, Ser CT, et al. Autonomous chemical experiments: Challenges and perspectives on establishing a self-driving lab. *Acc Chem Res.* 2022;55(17):2454-66.
- Baird SG, Sparks TD. What is a minimal working example for a self-driving laboratory? *Matter.* 2022;5(12):4170-8.
- Zhou ZH. Machine learning. Singapore: Springer Nature; 2021.
- Roch LM, Häse F, Kreisbeck C, Tamayo-Mendoza T, Yunker LP, Hein JE, et al. ChemOS: Orchestrating autonomous experimentation. *Sci Robot.* 2018;3(19):eaat5559.
- Kavalsky L, Hegde VI, Muckley E, Johnson MS, Meredig B, Viswanathan V. By how much can closed-loop frameworks accelerate computational materials discovery? *Digit Discov.* 2023;2(4):1112-25.
- He T. Machine learning inorganic solid-state synthesis from materials science literature. Berkeley (CA): University of California, Berkeley; 2023.

Noé F, Tkatchenko A, Müller KR, Clementi C. Machine learning for molecular simulation. *Annu Rev Phys Chem.* 2020;71(1):361-90.

Aubin CA, Buskohl PR, Vaia RA, Shepherd RF. Autonomous material systems: CA Aubin et al. *MRS Bull.* 2024;49(10):1070-8.

MacLeod BP, Parlange FG, Rupnow CC, Dettelbach KE, Elliott MS, Morrissey TD, et al. A self-driving laboratory advances the Pareto front for material properties. *Nat Commun.* 2022;13(1):995.

Dai Y, Lu P, Cao Z, Campbell CT, Xia Y. The physical chemistry and materials science behind sinter-resistant catalysts. *Chem Soc Rev.* 2018;47(12):4314-31.

Bennett JA, Abolhasani M. Autonomous chemical science and engineering enabled by self-driving laboratories. *Curr Opin Chem Eng.* 2022;36:100831.

Catania F, de Souza Oliveira H, Lugoda P, Cantarella G, Münzenrieder N. Thin-film electronics on active substrates: Review of materials, technologies and applications. *J Phys D Appl Phys.* 2022;55(32):323002.

Vriza A, Chan H, Xu J. Self-driving laboratory for polymer electronics. *Chem Mater.* 2023;35(8):3046-56.

Holland I, Davies JA. Automation in the life science research laboratory. *Front Bioeng Biotechnol.* 2020;8:571777.

Wigh DS, Goodman JM, Lapkin AA. A review of molecular representation in the age of machine learning. *Wiley Interdiscip Rev Comput Mol Sci.* 2022;12(5):e1603.

Li J, Zhou M, Wu HH, Wang L, Zhang J, Wu N, et al. Machine learning-assisted property prediction of solid-state electrolyte. *Adv Energy Mater.* 2024;14(20):2304480.

Ferguson AL, Brown KA. Data-driven design and autonomous experimentation in soft and biological materials engineering. *Annu Rev Chem Biomol Eng.* 2022;13(1):25-44.

Oganov AR, Pickard CJ, Zhu Q, Needs RJ. Structure prediction drives materials discovery. *Nat Rev Mater.* 2019;4(5):331-48.