

ORIGINAL RESEARCH

Open access

The Problem of Scientific Regret in AI-Driven Materials Selection

Hiroshi Tanaka¹, Yuki Sato^{1*}, Kenji Mori², Rina Okabe¹, Takashi Ito²

Abstract

The accelerating integration of artificial intelligence into materials selection processes has brought unprecedented efficiency to high-throughput screening and discovery campaigns, yet it has also introduced a subtle but profound failure mode that remains largely unrecognized in the field: scientific regret. This paper identifies scientific regret as a distinct failure mode in AI-driven materials science—the ex-post realization that a better material or research direction was passed over due to an AI recommendation, often under conditions of irreducible uncertainty and vast combinatorial search spaces. Unlike traditional statistical errors, scientific regret captures the experiential and consequential dimension of missed opportunities in research trajectories that are difficult or impossible to reverse. Drawing on foundational work in decision theory and recent advances in Bayesian optimization for materials discovery, the paper defines scientific regret, delineates its mechanisms of production within AI systems, develops a typology tailored to materials contexts, and outlines principles for its detection and mitigation. By analyzing how premature search space pruning, overconfidence in negative predictions, and misaligned acquisition functions contribute to regret, this analysis reveals how current AI paradigms may systematically undervalue exploration in favor of short-term gains. The implications for materials AI practice are significant, calling for the design of regret-sensitive systems that better balance exploitation with the long-term costs of locked-in choices. Ultimately, embracing scientific regret as a core design constraint promises to foster more robust, reflective, and innovative approaches to autonomous materials research. Scientific regret is not merely an abstract philosophical concern but a practical barrier to genuine progress in materials science. When AI systems guide researchers away from promising chemistries or structures, the subsequent realization of a missed opportunity can stall entire research programs, waste limited experimental resources, and distort the collective knowledge base of the field. This failure mode is especially insidious because materials discovery operates in enormous design spaces where exhaustive enumeration is impossible and where negative predictions are rarely revisited once resources are committed elsewhere. By foregrounding scientific regret as a failure mode, this analysis seeks to reorient the community toward decision frameworks that explicitly account for the irreversible nature of many AI-influenced choices in materials selection.

Keywords Bayesian optimization, Counterfactual reasoning, Scientific regret, AI-driven materials selection, Decision theory under uncertainty, Opportunity cost

*Correspondence:

Yuki Sato

yuki.sato@outlook.com

¹ Department of Intelligent Materials Engineering, University of Tokyo, Tokyo, Japan

² Department of AI-Based Materials Discovery, Kyoto University, Kyoto, Japan

Introduction

The rapid adoption of artificial intelligence techniques in materials science has transformed the process of materials selection and discovery. Machine learning models now routinely guide researchers toward promising compositions, structures, and processing conditions by predicting properties from vast chemical spaces. Systems employing Bayesian optimization and active learning have demonstrated remarkable efficiency in navigating high-dimensional design spaces

that would be intractable for traditional Edisonian approaches [1–5]. However, these AI-driven recommendations often involve decisions that are effectively irreversible or carry substantial costs in terms of time, resources, and scientific opportunity. When an AI system recommends against pursuing a particular material candidate or research direction, and that choice later proves suboptimal, researchers may experience what this paper terms scientific regret.

Scientific regret arises in the context of AI-assisted decision-making where the consequences of forgoing an alternative pathway become apparent only after significant investment has been made along the recommended path. This phenomenon draws directly from foundational work in decision theory under uncertainty [1–3], where regret is understood as the emotional and cognitive response to realizing that a different choice would have led to a superior outcome. In materials science, such regret is particularly acute because the search space is combinatorially enormous, and experimental validation remains expensive and time-consuming [6, 7]. Once a promising lead is deprioritized based on AI predictions, it may never be revisited due to path dependency in research programs and funding constraints.

Despite the proliferation of literature on machine learning applications in materials informatics [8–10], the concept of scientific regret has received little systematic attention. Most studies focus on metrics such as prediction accuracy, sample efficiency, or discovery rates, without addressing the downstream costs of erroneous negative recommendations or overly narrow exploration strategies [8, 11–14]. This oversight is significant because materials selection decisions frequently lock in subsequent research trajectories. For instance, committing to a particular class of perovskites [6] based on early AI screening may preclude investigation of entirely different chemistries that could offer better stability or performance under real-world conditions, as seen in electrocatalyst discovery campaigns [5].

The problem is compounded by the nature of scientific publication, which tends to favor positive results and successful discoveries while underreporting negative outcomes or abandoned directions. Consequently, the true extent of scientific regret in the field remains hidden. This paper addresses this gap by conducting a failure mode analysis of scientific regret in AI-driven materials selection. It begins by formally defining the concept and distinguishing it from related notions such as opportunity cost and Type II errors. Subsequent sections examine why regret carries particular weight in materials science, analyze the mechanisms by which AI systems generate regret, and develop a typology of regret types specific to this domain. The analysis then turns to principles for detecting and mitigating regret before exploring its relationship to other failure modes and implications for practice.

By theorizing scientific regret, this study aims to shift the discourse in materials AI from a narrow focus on optimization performance toward a more holistic consideration of long-term scientific value and decision quality under uncertainty. The integration of regret considerations into AI design promises not only to reduce wasted effort but also to preserve the serendipitous potential that has historically driven breakthroughs in materials discovery. In an era where autonomous research platforms are increasingly proposed [7], acknowledging scientific regret becomes essential to ensuring that AI augments rather than constrains human scientific creativity.

Defining Scientific Regret

Scientific regret represents a specific form of decision regret tailored to the scientific enterprise, particularly when amplified by AI-assisted tools in fields like materials science. It captures the retrospective assessment that an earlier choice, influenced by algorithmic recommendation, resulted in the abandonment of a more fruitful path.

Scientific regret is the ex-post realization by a research team or individual scientist that an AI-driven decision in materials selection or experimental prioritization led to the forgoing of a superior material, property, or research trajectory, where the forgone alternative would have produced demonstrably better scientific or technological outcomes, had the resources been allocated differently.

This definition emphasizes several distinguishing features. First, it is inherently retrospective and counterfactual in nature, aligning with foundational theories of regret in economics and psychology [2, 3]. Unlike mere opportunity cost, which is an

ex-ante calculation of foregone benefits at the moment of decision [1], scientific regret manifests after evidence emerges suggesting the choice would have been preferable. It differs from a simple type II error (false negative) in high-throughput screening because it incorporates the emotional, cognitive, and institutional dimensions of scientific decision-making under uncertainty, not merely statistical misclassification. While a type II error describes the probability of missing a positive instance, scientific regret concerns the lived experience and long-term impact of having acted upon that error within a path-dependent research program.

Furthermore, scientific regret is distinct from general counterfactual loss because it specifically arises within the context of AI-mediated choices that constrain the exploration of vast chemical spaces. In materials discovery, the decision to deprioritize a candidate based on model predictions [5, 13] often carries downstream consequences that amplify the sense of regret when later discoveries validate the missed opportunity. The definition thus positions scientific regret as a distinct failure mode that bridges decision theory [1] with the practical realities of autonomous or semi-autonomous materials research [7].

Elaborating further, scientific regret can be conceptualized as the difference between the utility of the chosen path and the utility of the best forgone alternative, weighted by the probability or evidence that emerges later. It is not simply the absence of optimality but the recognition that AI recommendations contributed to suboptimal path selection in ways that were avoidable with a different algorithmic design. By formalizing this concept, the paper provides a lens through which to critique current practices in machine learning for materials informatics [9, 10] and to propose more regret-aware methodologies. The retrospective nature also highlights why regret is especially salient in materials science: experimental synthesis and characterization are resource-intensive, and once a direction is abandoned, revisiting it requires new funding cycles or shifts in research priorities that are rarely forthcoming. This temporal asymmetry distinguishes scientific regret from static error metrics and underscores its status as a dynamic failure mode that evolves over the lifecycle of a research project.

Table 1 delineates scientific regret from adjacent constructs, clarifying its unique role as a consequential, path-dependent failure mode rather than a purely statistical or theoretical artifact.

Table 1. Conceptual differentiation of scientific regret from adjacent failure modes in AI-driven materials selection

Dimension	Scientific regret	Type II error (false negative)	Opportunity cost	Path dependence	Model overconfidence
Temporal orientation	Ex-post (retrospective realization)	Ex-ante probabilistic	Ex-ante evaluation	Dynamic over time	Ex-ante prediction
Core definition	Realized loss from AI-guided forgone superior pathway	Missed detection of true positive	Value of next-best alternative at decision time	Lock-in from early choices	Miscalibrated certainty in predictions
Domain scope	AI-mediated scientific workflows	Statistical classification	General decision theory	Process dynamics in research trajectories	Model behavior
Human/Institutional impact	High (affects research direction, funding, and publications)	Low (primarily technical metric)	Moderate (theoretical planning tool)	High (structural constraint)	Indirect (via decisions)

Reversibility	Typically low (path-dependent and resource-constrained)	High (can re-test)	Not applicable	Very low	Moderate (via recalibration)
Visibility in literature	Largely hidden	Explicitly measured	Explicitly modeled	Recognized but diffuse	Widely studied
Relation to AI systems	Emergent failure mode of AI-driven decisions	Output error	Decision framing concept	Emergent process outcome	Model property
Primary consequence	Lost scientific opportunity	Reduced accuracy	Suboptimal allocation	Entrenched trajectories	Misleading guidance
Mitigation strategy	Regret-aware design principles	Improve recall	Better planning	Increase flexibility	Calibration techniques

In summary, scientific regret is not a mere statistical artifact but a failure mode rooted in the interplay between algorithmic guidance, human judgment, and the irreversible commitments inherent to materials experimentation. Its formal definition serves as the foundation for the subsequent analysis of why it matters, how it is produced, and how it might be mitigated.

Why Regret Matters in Materials Science

Scientific regret carries particular weight in materials science for three interlocking reasons that arise directly from the unique characteristics of the field. First, materials selection decisions are frequently irreversible or carry prohibitive reversal costs. Once an AI system recommends deprioritizing a candidate chemistry, the associated experimental resources—synthesis equipment time, characterization facilities, and personnel effort—are redirected elsewhere, making later reconsideration logistically and financially burdensome [7]. In active learning campaigns for electrocatalysts, for example, early negative predictions can terminate entire lines of inquiry that later prove viable under slightly different conditions [5].

The search spaces in materials discovery are astronomically large, rendering passed-over materials effectively lost forever unless explicit mechanisms for revisitation are built into the workflow. High-throughput virtual screening and Bayesian optimization routinely evaluate millions of candidates, yet only a tiny fraction receive experimental follow-up [8, 11]. When AI narrows the explored subspace too aggressively, promising materials outside the selected region may remain undiscovered indefinitely, compounding regret at the scale of entire subfields. This vastness distinguishes materials science from domains where alternatives can be revisited with minimal incremental cost.

Publication and funding biases against negative results and abandoned directions further amplify scientific regret. The scientific community rarely learns about materials that were dismissed by AI recommendations, creating a systematic underestimation of false negatives and perpetuating the very decision biases that generated the regret in the first place [9, 13]. This feedback loop means that regret not only affects individual research teams but also distorts collective progress across the discipline.

Taken together, these factors make scientific regret a uniquely consequential failure mode. It is not simply an individual researcher’s disappointment but a structural inefficiency that undermines the very promise of AI-accelerated discovery. By ignoring regret, current materials AI practices risk optimizing for short-term publication metrics while sacrificing the long-term robustness and serendipity that have historically driven breakthroughs in the field.

Mechanisms of Regret Production

AI systems generate scientific regret through three primary mechanisms that interact within typical materials discovery pipelines.

Premature search space pruning occurs when algorithms eliminate large regions of chemical or structural space based on early, low-confidence predictions. In Bayesian optimization frameworks, acquisition functions often favor exploitation over broad exploration, leading to rapid contraction of the considered space [4, 8]. Once pruned, these regions are rarely re-evaluated, locking researchers into narrower trajectories that later prove suboptimal. This mechanism is particularly problematic in materials science because initial surrogate models trained on limited data can confidently dismiss entire families of compounds that would have succeeded under refined conditions.

Overconfidence in negative predictions arises when machine learning models assign excessively high certainty to “no-go” classifications without adequate uncertainty quantification. Modern materials informatics pipelines frequently present false negatives as definitive, encouraging researchers to abandon candidates that, upon deeper investigation or with additional data, reveal hidden potential [13, 14]. The mechanism is exacerbated by the black-box nature of many deep learning architectures used for property prediction, which obscures the epistemic uncertainty inherent in extrapolating to novel chemistries.

Misaligned acquisition functions prioritize immediate expected improvement while undervaluing long-term discovery potential. Standard formulations in Bayesian optimization focus on short-horizon gains, systematically favoring incremental refinements over high-risk, high-reward explorations that could uncover transformative materials [4, 15–17]. In practice, this misalignment directs resources toward incrementally better but ultimately mediocre candidates while sidelining truly novel directions.

These mechanisms are not isolated but reinforce one another, creating feedback loops that entrench suboptimal paths in materials AI workflows.

Figure 1 presents a unified process model illustrating how AI decision inputs generate scientific regret through identifiable mechanisms, observable detection signatures, and targeted mitigation interventions.

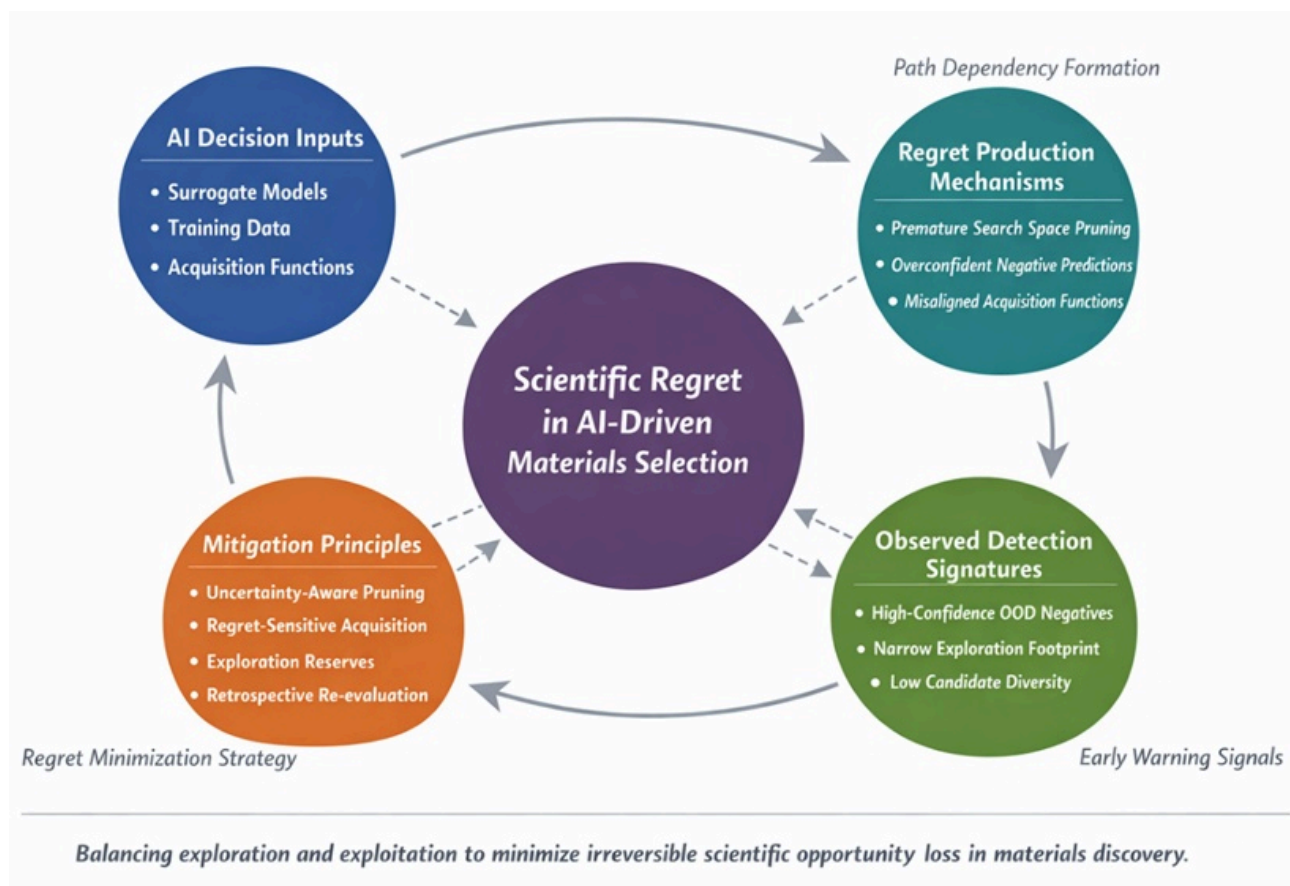


Figure 1. A unified process model illustrating how AI decision inputs generate scientific regret through identifiable mechanisms, observable detection signatures, and targeted mitigation interventions

A Typology of Scientific Regret

A typology of scientific regret clarifies its varied manifestations in materials science contexts.

False negative regret occurs when an AI system incorrectly classifies a promising material as unviable, leading researchers to discard it outright. This type is common in high-throughput screening, where models trained on limited data produce overly pessimistic predictions for out-of-distribution candidates [5, 6]. The regret materializes when independent validation or serendipitous later work reveals the material's superior properties.

Path dependency regret emerges from early AI recommendations that lock research programs into trajectories from which escape becomes costly. Initial choices of surrogate models or acquisition strategies can create self-reinforcing loops that exclude alternative chemistries, even as new data suggest better routes [7, 11]. This regret is cumulative and particularly difficult to reverse once publications, patents, and funding are tied to the chosen path.

Resource allocation regret arises when AI directs finite experimental resources toward low-value targets while higher-value opportunities wait indefinitely. Misguided prioritization based on short-term acquisition functions can starve promising but initially uncertain directions of the data needed for accurate evaluation [8, 18].

The typology highlights that regret is not monolithic but appears in distinct forms, each requiring tailored analytical and design responses. By distinguishing these types, researchers can better anticipate where regret is likely to arise and intervene proactively.

Detection Principles

Detection of scientific regret in AI-driven materials selection requires systematic attention to early conceptual signatures that indicate when an AI system is likely generating avoidable missed opportunities. Rather than relying on post-hoc empirical validation, which is often infeasible once resources have been committed, detection principles focus on observable patterns within the decision-making workflow itself. One critical signature is the presence of high-confidence negative predictions on novel chemistries or structures that lie far from the training distribution. When surrogate models assign near-certainty to “no-go” classifications for candidates that differ substantially in composition or bonding from previously evaluated materials, this frequently foreshadows false negative regret [13, 14]. Such overconfidence is particularly problematic in materials informatics because chemical space is vast and sparsely sampled; a model trained primarily on oxide perovskites, for instance, may dismiss entirely new classes of halide or sulfide candidates with unwarranted certainty, even though subsequent independent studies may reveal superior properties [5, 6].

Another key detection signature emerges as narrow exploration relative to the overall search space. In Bayesian optimization loops, this appears when the acquisition function rapidly contracts the considered subspace after only a limited number of iterations, leaving large regions of potential materials unexamined [4, 8, 11]. Monitoring the ratio of explored volume to total combinatorial space, or tracking the diversity of sampled candidates across successive batches, can reveal this pattern before irreversible experimental commitments are made. When the algorithm consistently returns candidates clustered within a small neighborhood of previously successful leads, the risk of path dependency regret increases sharply because promising but initially uncertain directions are effectively excluded from consideration [7, 18].

A further signature is reflected in a persistent lack of diversity in recommended candidates, signaling resource allocation regret. When the AI pipeline repeatedly proposes incremental variants of the same structural family or processing route while sidelining higher-risk, higher-reward alternatives, the system implicitly directs finite experimental budgets toward low-variance improvements at the expense of transformative discovery [17, 18]. This condition can be detected using metrics such as Tanimoto similarity across successive recommendation sets or by examining the entropy of the proposed candidate pool; persistently low diversity values are strongly associated with later recognition that superior materials were overlooked [9, 10].

Taken together, these signatures—high-confidence negatives on out-of-distribution candidates, narrow exploration footprints, and low candidate diversity—function as practical sentinels within live workflows. They enable research teams to identify potential regret before synthesis or characterization resources are fully expended. A useful conceptual aid is a decision-tree representation in which each AI-mediated choice point is modeled as a node. At each node, the selected branch is pursued. In contrast, alternative branches are annotated with an estimated regret magnitude, representing the prospective gap between the realized outcome and the best forgone alternative. The width of each pruned branch can be scaled according to the confidence of the negative prediction and the remaining unexplored volume, providing an intuitive visualization of how regret accumulates across multiple decision layers [1–3].

By embedding these detection principles directly into the materials AI pipeline, practitioners can shift from reactive regret assessment to proactive identification. These principles do not require additional experiments or datasets; instead, they rely on introspection of model uncertainty, the geometry of explored space, and the statistical properties of recommendation sets. When applied consistently, they transform scientific regret from an invisible downstream cost into a monitorable design parameter, enabling timely intervention before path-dependent lock-in occurs [7, 11]. In this way, detection becomes an integral component of responsible AI-assisted materials selection rather than a retrospective consideration.

Mitigation Principles

Mitigation of scientific regret requires a deliberate redesign of AI systems so that they explicitly account for the cost of missed opportunities, rather than optimizing solely for immediate expected improvement. This shift can be achieved

through a set of complementary design principles that collectively reconfigure how candidates are evaluated, retained, and revisited throughout the discovery process.

A central adjustment involves moving away from rigid elimination thresholds toward uncertainty-aware pruning strategies. Conventional pipelines frequently discard candidates whose predicted properties fall below fixed cutoffs, even when the associated uncertainty remains substantial [4, 8]. Such practices implicitly treat predictions as definitive, leading to premature exclusion of potentially valuable materials. In contrast, retaining candidates whose uncertainty intervals overlap with high-performing regions allows the system to defer irreversible decisions until additional evidence is available. This is particularly important in materials discovery, where sparse data and vast chemical space make early overconfidence especially costly [13, 14]. Operationally, this approach can be implemented through dynamic retention mechanisms—such as maintaining a revisitable pool of uncertain candidates that are periodically reassessed as new data accumulates.

A related modification concerns the objective functions guiding candidate selection. Standard acquisition strategies in Bayesian optimization prioritize expected improvement or confidence bounds without accounting for the opportunity cost of excluded alternatives. Incorporating regret-sensitive adjustments into these functions introduces a counterbalancing force, penalizing decisions that prematurely eliminate large or uncertain regions of the search space [2, 3]. By weighting this penalty according to both uncertainty and the size of the unexplored domain, the system is encouraged to preserve exploratory pathways that might otherwise be abandoned. This aligns short-term optimization behavior with longer-term scientific value, reducing the tendency toward overly conservative, incremental discovery trajectories [4, 17].

Beyond algorithmic adjustments, structural safeguards can be embedded at the level of resource allocation. Allocating a dedicated portion of the experimental budget to exploratory sampling—particularly in underrepresented or low-confidence regions—ensures that diversity is maintained even when the dominant optimization logic favors exploitation [11, 18-22]. This mechanism acts as a buffer against path dependency, preserving alternative discovery trajectories that may later prove more fruitful. By maintaining parallel lines of inquiry, it becomes possible to adapt more effectively when initial assumptions or model predictions are revised [7, 10].

Finally, the integration of systematic retrospective analysis introduces an additional layer of resilience. Periodic re-evaluation of previously discarded candidates using updated models and newly available data allows the system to recover opportunities that may have been overlooked earlier [7, 9]. Rather than treating pruning decisions as final, this approach reframes them as provisional, subject to revision as knowledge evolves. In doing so, it transforms scientific regret from an irreversible loss into a detectable and, in some cases, correctable signal.

Taken together, these design principles reorient materials AI from a narrow focus on efficiency toward a more balanced emphasis on robustness and long-term discovery potential. They acknowledge the irreversible nature of many materials decisions and the structural uncertainty inherent in exploring vast chemical spaces [5, 6, 8]. Importantly, their implementation does not require fundamental changes to existing frameworks; instead, they can be incorporated through targeted modifications to selection criteria, budgeting strategies, and evaluation protocols. The result is a class of AI systems that not only accelerate discovery but also actively guard against the systematic loss of promising alternatives, thereby supporting a more creative and resilient scientific process. Table 2 integrates mechanisms, detection signatures, and mitigation strategies into a unified framework, enabling systematic intervention across the AI-driven materials discovery pipeline.

Table 2. Mechanism– detection–mitigation alignment framework for scientific regret in materials AI

Regret production mechanism	Operational description	Observable detection signature	Analytical indicator	Targeted mitigation principle	Expected system-level effect

Premature search space pruning	Early elimination of large candidate regions under limited data	Rapid contraction of explored space	Low explored volume ratio; clustering of candidates	Uncertainty-aware pruning	Preserves optionality and reduces irreversible exclusion
Overconfidence in negative predictions	Excessive certainty in “no-go” classifications	High-confidence negatives on OOD candidates	Narrow uncertainty intervals; calibration error	Retrospective analysis protocols	Enables recovery of falsely discarded candidates
Misaligned acquisition functions	Optimization favors short-term gains over long-term discovery	Repeated selection of incremental variants	Low diversity (entropy); high similarity scores	Regret-sensitive acquisition functions	Rebalances the exploration–exploitation trade-off
Resource allocation bias	Concentration of experimental effort on low-risk options	Persistent low diversity in recommendation sets	Tanimoto similarity; diversity entropy metrics	Exploration reserves	Ensures parallel exploration of high-risk candidates
Path dependency reinforcement	Feedback loop narrows future search directions	Reduced variability across iterations	Decreasing search entropy over time	Combined mitigation (all four principles)	Maintains multiple viable research trajectories

Relationship to Other Failure Modes

Scientific regret does not exist in isolation; it intersects with several well-documented failure modes in materials AI while remaining conceptually distinct. Overconfidence, for example, is a primary driver of the false negatives that later generate regret, yet the two phenomena differ in focus. Overconfidence concerns the statistical reliability of model outputs—specifically the tendency of deep learning or Gaussian process models to assign narrow credible intervals even when extrapolating to novel chemistries [13, 14]. Regret, by contrast, emphasizes the downstream experiential and institutional cost of acting on those overconfident negatives: the abandoned synthesis route, the redirected funding, and the missed publication opportunity. Fixing overconfidence through better calibration or ensemble methods may reduce the incidence of false negatives. Still, it does not automatically eliminate regret unless the system is also redesigned to preserve uncertain candidates and revisit discarded options [7, 11, 23–25].

Path dependence is another closely related failure mode. It describes the self-reinforcing lock-in that occurs when early AI choices shape subsequent data collection and model retraining, progressively narrowing the feasible research trajectory [7, 11]. While path dependence explains the mechanism by which early decisions become difficult to reverse, scientific regret captures the later realization of the scientific price paid for that lock-in. In other words, path dependence is the process; regret is the evaluative judgment that the process led to a suboptimal outcome. The typology developed in Section 5 clarifies this distinction by separating path dependency regret as one specific manifestation among others.

Representation bias in training data—where models reflect the limited diversity of existing literature—can also feed into regret by systematically undervaluing underrepresented material classes [9, 10]. Yet representation bias is a data-generation problem, whereas regret is a decision-outcome problem. Even a perfectly representative dataset can still produce regret if the acquisition function or pruning logic fails to balance exploration adequately. The relationship is therefore contributory rather than identical: representation bias increases the baseline probability of regret, but regret can arise independently through misaligned decision rules even in unbiased datasets.

By mapping these relationships explicitly, the framework avoids conflating distinct concepts. Overconfidence, path dependence, and representation bias are root causes or enabling conditions, while scientific regret is the consequential failure mode that ultimately matters to researchers and funding agencies. Recognizing the distinctions allows targeted interventions: calibration techniques address overconfidence, diversity-aware sampling counters representation bias, and the mitigation principles in Section 7 directly tackle regret itself. This nuanced mapping strengthens the overall failure-mode analysis and prevents the common error of assuming that solving one problem automatically solves the others [1, 2, 8, 26–29].

Implications for Materials AI Practice

Recognizing scientific regret as a core failure mode carries several important implications for how the materials AI community designs, evaluates, and deploys its systems. One immediate consequence is the need to expand reporting practices beyond conventional success metrics. Current studies tend to emphasize successful discoveries and sample efficiency, while providing little visibility into the candidates that were confidently rejected [5, 6]. Incorporating systematic reporting of false negative rates and uncertainty footprints—particularly for candidates outside the training manifold—would expose otherwise hidden forms of regret and enable collective learning from missed opportunities rather than allowing them to remain unexamined [13, 14].

A related implication concerns the structure of discovery workflows themselves. Abandoned directions should not be treated as irrecoverable endpoints but instead preserved in a form that allows future reconsideration. This can be achieved through persistent candidate repositories that retain pruned materials alongside their associated uncertainty estimates, as well as through automated re-ranking procedures that revisit discarded options when new data becomes available [7, 8]. By keeping alternative pathways accessible, such mechanisms counteract path dependency and ensure that missed opportunities can be revisited as knowledge evolves.

The role of regret also extends to the core design criteria of AI systems. Rather than focusing exclusively on predictive accuracy or expected improvement, model selection and optimization strategies should be evaluated in terms of their expected regret under realistic search conditions [4, 17]. This introduces a decision-theoretic perspective that prioritizes robustness, long-term exploration, and openness to unexpected discoveries over short-term performance gains [9, 10].

Taken together, these implications promote a more reflective and responsible approach to AI-assisted materials selection. They do not hinder the pace of discovery; instead, they enhance its sustainability by preserving optionality within vast design spaces and by ensuring that the intellectual and material investments of research yield a more complete and enduring scientific return.

Conclusion

Scientific regret represents a distinct and previously under-theorized failure mode in AI-driven materials selection: the ex-post realization that a superior material or research direction was passed over because of algorithmic recommendation under conditions of irreducible uncertainty. This paper has defined the concept, distinguished it from related notions such as opportunity cost and type II error, analyzed the mechanisms that produce it, developed a typology of its manifestations, outlined detection signatures, and proposed mitigation principles grounded in decision theory and materials informatics. By foregrounding regret as a design constraint rather than an inevitable byproduct of efficiency, the field can move toward AI systems that better balance short-term gains with long-term scientific openness.

The ultimate goal is not merely faster discovery but more robust, creative, and trustworthy autonomous materials research. Embracing regret-aware design promises to reduce wasted resources, preserve serendipitous potential, and align algorithmic guidance more closely with the irreversible realities of experimental science. Future work in materials AI

should therefore treat scientific regret as a first-class consideration, ensuring that the powerful tools now at our disposal expand rather than constrain the horizons of possibility.

Acknowledgements

None

Conflict of interest

None

Financial Support

None

Ethics statement

None

Received: 10 Jul 2021

Revised: 24 Aug 2021

Accepted: 09 Oct 2021

Published online: 18 January 2022

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

- Bartholomeyczik K, Gusenbauer M, Treffers T. The influence of incidental emotions on decision-making under risk and uncertainty: A systematic review and meta-analysis of experimental evidence. *Cogn Emot.* 2022;36(6):1054-73.
<https://doi.org/10.1080/02699931.2022.2099349>.
- Schünemann HJ, Reinap M, Piggott T, Laidmäe E, Köhler K, Pöld M, et al. The ecosystem of health decision making: From fragmentation to synergy. *Lancet Public Health.* 2022;7(4):378-90.
[https://doi.org/10.1016/S2468-2667\(22\)00057-3](https://doi.org/10.1016/S2468-2667(22)00057-3).
- Zhang T, Zhang Q, Wu J, Wang M, Li W, Yan J, et al. The critical role of the orbitofrontal cortex for regret in an economic decision-making task. *Brain Struct Funct.* 2022;227(8):2751-67.
<https://doi.org/10.1007/s00429-022-02568-w>.
- Tran-The H, Gupta S, Rana S, Venkatesh S. Regret bounds for expected improvement algorithms in Gaussian process bandit optimization. In: *Proceedings of the 25th International Conference on Artificial Intelligence and Statistics AISTATS*; 2022. p. 8715-

37.

Tran K, Ulissi ZW. Active learning across intermetallics to guide discovery of electrocatalysts for CO₂ reduction and H₂ evolution. *Nat Catal.* 2018;1(9):696-703.

Assirey EA. Perovskite synthesis, properties and their related biochemical and industrial application. *Saudi Pharm J.* 2019;27(6):817-29.

National Research Council, Division on Engineering, Physical Sciences, Board on Physics, Physics Survey Overview Committee. Physics in a new era: An overview.
<https://doi.org/10.17226/10118>.

Liang Q, Gongora AE, Ren Z, Tiihonen A, Liu Z, Sun S, et al. Benchmarking the performance of Bayesian optimization across multiple experimental materials science domains. *npj Comput Mater.* 2021;7(1):188.

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.

Ramprasad R, Batra R, Pilania G, Mannodi-Kanakkithodi A, Kim C. Machine learning in materials informatics: Recent applications and prospects. *npj Comput Mater.* 2017;3(1):54.

Lookman T, Balachandran PV, Xue D, Yuan R. Active learning in materials science with emphasis on adaptive sampling using uncertainties for targeted design. *npj Comput Mater.* 2019;5(1):21.

Peterson JC, Bourgin DD, Agrawal M, Reichman D, Griffiths TL. Using large-scale experiments and machine learning to discover theories of human decision-making. *Science.* 2021;372(6547):1209-14.

Butler KT, Davies DW, Cartwright H, Isayev O, Walsh A. Machine learning for molecular and materials science. *Nature.* 2018;559(7715):547-55.

Zhong X, Gallagher B, Liu S, Kaikhura B, Hiszpanski A, Han TY. Explainable machine learning in materials science. *npj Comput Mater.* 2022;8(1):204.

Goodall RE, Parackal AS, Faber FA, Armiento R, Lee AA. Rapid discovery of stable materials by coordinate-free coarse graining. *Sci Adv.* 2022;8(30):eabn4117.

Fare C, Fenner P, Benatan M, Varsi A, Pyzer-Knapp EO. A multi-fidelity machine learning approach to high throughput materials screening. *npj Comput Mater.* 2022;8(1):257.

Stuke A, Rinke P, Todorović M. Efficient hyperparameter tuning for kernel ridge regression with Bayesian optimization. *Mach Learn Sci Technol.* 2021;2(3):035022.

Balachandran PV. Adaptive machine learning for efficient materials design. *MRS Bull.* 2020;45(7):579-86.

Zunger A. Inverse design in search of materials with target functionalities. *Nat Rev Chem.* 2018;2(4):0121.

Omar ÖH, Del Cueto M, Nemataram T, Troisi A. High-throughput virtual screening for organic electronics: A comparative study of alternative strategies. *J Mater Chem C.* 2021;9(39):13557-83.

Shinde PP, Shah S. A review of machine learning and deep learning applications. In: 2018 Fourth International Conference on Computing, Communication, Control and Automation (ICCUBEA); 2018. p. 1-6.

Moosavi SM, Jablonka KM, Smit B. The role of machine learning in the understanding and design of materials. *J Am Chem Soc.* 2020;142(48):20273-87.

Zhou ZH. Machine learning. Singapore: Springer Nature; 2021. p. 20.

He Y, Cubuk ED, Allendorf MD, Reed EJ. Metallic metal-organic frameworks predicted by the combination of machine learning methods and ab initio calculations. *J Phys Chem Lett.* 2018;9(16):4562-9.

Ko TW, Finkler JA, Goedecker S, Behler J. General-purpose machine learning potentials capturing nonlocal charge transfer. *Acc Chem Res.* 2021;54(4):808-17.

Zhu G, Zhang H, Jacobson O, Wang Z, Chen H, Yang X, et al. Combinatorial screening of DNA aptamers for molecular imaging of HER2 in cancer. *Bioconjug Chem.* 2017;28(4):1068-75.

Ren F, Ward L, Williams T, Laws KJ, Wolverton C, Hattrick-Simpers J, et al. Accelerated discovery of metallic glasses through iteration of machine learning and high-throughput experiments. *Sci Adv.* 2018;4(4):eaq1566.

Jablonka KM, Jothiappan GM, Wang S, Smit B, Yoo B. Bias free multiobjective active learning for materials design and discovery. *Nat Commun.* 2021;12(1):2312.

Hestness J, Narang S, Ardalani N, Damos G, Jun H, Kianinejad H, et al. Deep learning scaling is predictable, empirically. *arXiv preprint arXiv:1712.00409.* 2017.
<https://doi.org/10.48550/arXiv.1712.00409>.