

ORIGINAL RESEARCH

Open access

The Problem of Scientific Lock-In Through Early AI Adoption in Materials Domains

Hassan Ali¹, Mariam Farooq^{1*}, Usman Shah²

Abstract

The rapid proliferation of artificial intelligence (AI) techniques across materials science domains has delivered unprecedented predictive power and design acceleration. Yet, it has simultaneously engendered a previously under-examined failure mode: scientific lock-in through early AI adoption in materials domains. Scientific lock-in is defined here as the self-reinforcing entrenchment of specific AI methods, representations, or frameworks chosen early in the development of a subfield, rendering subsequent adoption of demonstrably superior alternatives prohibitively difficult even when their advantages become evident to the community. This failure mode arises through four interlocking mechanisms—increasing returns, switching costs, network effects, and institutionalization—and manifests across four distinct types: representational, methodological, data, and evaluation lock-in, each of which is shown to constrain the epistemic possibilities of materials research in characteristic ways. The resultant failure modes include suboptimal persistence of inferior approaches, innovation suppression of promising alternatives, comparative ignorance that prevents fair benchmarking, and collective regret in which the community recognizes the problem yet remains collectively unable to escape it. Detection principles grounded in observable indicators such as method concentration, citation bias, switching resistance, and comparative gaps are proposed, while mitigation principles centered on methodological pluralism, standardized comparisons, modular interoperability, community audits, and targeted funding for alternatives offer practical pathways to preserve long-term adaptability. By framing scientific lock-in as a distinct failure mode in materials AI, the present analysis urges the community to treat early adoption choices not merely as technical decisions but as high-stakes commitments whose downstream consequences must be deliberately managed if the field is to retain its capacity for genuine scientific progress.

Keywords Materials informatics, Network effects, Path dependence, Scientific lock-in, Early AI adoption, Methodological pluralism

*Correspondence:

Mariam Farooq
mariam.farooq@outlook.com

¹ Department of AI-Based Materials Engineering, University of Lahore, Lahore, Pakistan

² Department of Smart Materials Analytics, National University of Sciences and Technology, Islamabad, Pakistan

Introduction

Materials artificial intelligence stands at a pivotal juncture. The field has witnessed an extraordinary acceleration in the development and deployment of machine-learning methods tailored to molecular and solid-state systems, with early landmark contributions establishing foundational pipelines for property prediction, inverse design, and high-throughput screening. These early successes, while undeniably transformative, have also seeded a subtle yet pervasive risk: the very speed and enthusiasm with which the community embraced particular AI methods, representations, and data infrastructures may now be locking research trajectories into paths that future innovations will find difficult to escape. Early adopters select specific graph-network architectures for crystal modeling or particular descriptor sets for featurization; these selections rapidly accumulate users, tools, citations, and institutional support; and the resulting

momentum renders alternative approaches—however theoretically superior—practically invisible or prohibitively costly to pursue [1-5].

The present paper, therefore, identifies scientific lock-in as a distinct failure mode in materials AI. Unlike transient inefficiencies or isolated modeling errors, scientific lock-in is a community-level, path-dependent process in which early technical choices become socially and institutionally self-sustaining, systematically constraining the future direction of inquiry [1]. This failure mode is especially acute in materials domains because the objects of study—complex, multiscale, data-scarce systems—amplify the consequences of any initial representational or algorithmic commitment. Once a particular featurization scheme or benchmark dataset becomes the de facto standard, entire subfields calibrate their expectations, validation protocols, and even funding proposals around it, creating a feedback loop that rewards conformity and penalizes deviation [4-10].

The analysis proceeds by first formalizing the concept of scientific lock-in and distinguishing it from related constructs such as general path dependence and technical lock-in. It then dissects the four core mechanisms that drive lock-in, articulates four primary types observable in materials AI, and develops a typology of the specific failure modes that lock-in produces. Throughout, the discussion remains conceptual and forward-looking, identifying risks and proposing safeguards rather than claiming retrospective empirical proof of lock-in in any single subfield. In doing so, the paper contributes a structured failure-mode framework that can inform both individual research practice and collective governance of the materials AI ecosystem.

Figure 1 presents the hierarchical architecture of scientific lock-in in materials AI, showing how early adoption choices cascade through self-reinforcing mechanisms into distinct lock-in types, downstream failure modes, and corresponding detection and mitigation pathways.

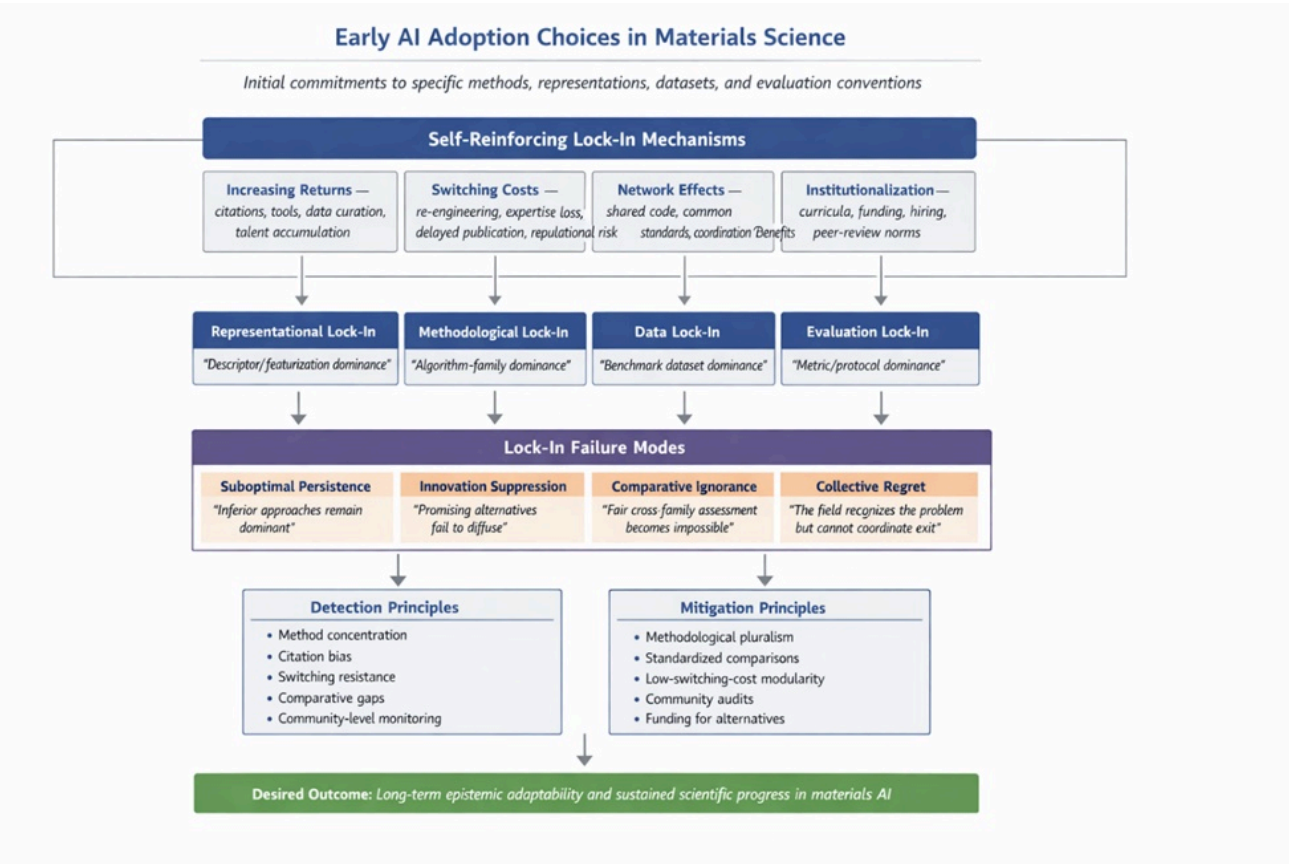


Figure 1. The hierarchical architecture of scientific lock-in in materials AI.

By treating early adoption decisions as high-leverage interventions whose long-term consequences must be anticipated, the community can avoid the regret that has afflicted other technology-intensive scientific domains where initial choices proved unexpectedly durable [11–16].

The urgency of this analysis stems from the current trajectory of the field. Recent reviews document how a handful of early methodological paradigms have come to dominate publications, toolkits, and curricula [2, 3]. First-mover advantages documented in broader AI-adoption studies [15–27] appear to operate with particular force in materials science, where the cost of generating new high-fidelity data is high, and the incentive to reuse established pipelines is correspondingly strong. Without deliberate countermeasures, the field risks trading short-term productivity gains for long-term epistemic rigidity. This introduction, therefore, sets the stage for a systematic examination of scientific lock-in, its mechanisms, manifestations, and remedies, with the explicit goal of equipping materials AI researchers and institutions with the conceptual tools needed to preserve methodological openness in an era of accelerating automation.

Defining Scientific Lock-In

Scientific lock-in constitutes a community-level failure mode in which an early-adopted AI method, representation, or framework acquires such dominance within a scientific subfield that switching to superior alternatives becomes effectively impossible despite clear evidence of their advantages. More formally:

Definition 1: Scientific lock-in — A situation in which an early AI method, representation, or framework becomes dominant in a research community, making it difficult or impossible to switch to superior alternatives despite recognized advantages.

This definition emphasizes three essential features: (i) the temporal precedence of the initial choice, (ii) the self-reinforcing dominance that follows, and (iii) the persistence of that dominance even after the community collectively acknowledges better options. Scientific lock-in is therefore distinct from mere path dependence, which describes any historical contingency that shapes future outcomes without necessarily implying entrapment or sub-optimality [11, 12]. Path dependence may be benign or even beneficial when early choices prove robust; scientific lock-in, by contrast, is pathological precisely because it locks the community into demonstrably inferior trajectories.

Scientific lock-in must also be differentiated from technical lock-in at the software or hardware level. Technical lock-in occurs when specific code libraries, file formats, or hardware architectures create compatibility barriers [16]. While technical lock-in can contribute to scientific lock-in, the latter operates at the level of collective scientific practice—citation patterns, peer-review norms, curriculum design, and funding priorities—rather than purely infrastructural constraints. A researcher may be able to run alternative code on the same hardware yet still face insurmountable social and epistemic barriers to doing so within the materials AI community.

Finally, scientific lock-in differs from cognitive lock-in at the individual researcher level. Cognitive lock-in refers to personal habits or mental models that make an individual resistant to new ideas [8]. Scientific lock-in, however, is an emergent property of the community: even researchers who privately recognize the limitations of the dominant approach may continue to employ it because the cost of deviating—lost citations, reviewer skepticism, or funding ineligibility—is borne individually. At the same time, the benefits of conformity are shared [20]. Thus, scientific lock-in is irreducibly social and structural, arising from the interplay of individual choices within an interconnected scientific ecosystem.

Table 1 clarifies the analytical boundaries of scientific lock-in by distinguishing it from adjacent forms of constraint that may coexist with, but do not fully explain, the community-level entrenchment described here.

Table 1. Analytical boundaries of scientific lock-in: distinguishing the construct from adjacent forms of constraint

Construct	Unit of analysis	Core defining feature	Trigger or source	What becomes constrained	Why it matters analytically for materials AI
Scientific lock-in	Research community/subfield	Early AI choices become self-reinforcing and persist despite recognized superior alternatives	Early adoption followed by increasing returns, switching costs, network effects, and institutionalization	Future methodological choice, comparative openness, and epistemic flexibility	Identifies a distinct community-level failure mode in which the field becomes trapped in inferior trajectories rather than merely shaped by history
Path dependence	Historical trajectory of a field or system	Earlier events shape later outcomes, but not necessarily inefficiently	Sequencing effects and contingent historical events	Direction of later development	Useful background condition, but too broad because it does not require entrapment or suboptimal persistence
Technical lock-in	Infrastructure, software, hardware, file formats	Compatibility barriers make switching difficult	Proprietary or sticky technical systems	Tools, workflows, interoperability	Helps explain infrastructural barriers, but does not capture how scientific legitimacy and peer norms reproduce dominance
Cognitive lock-in	Individual researcher	Personal habits or entrenched mental models resist new ideas	Familiarity, bounded rationality, prior training	Individual openness to alternatives	Important micro-foundation, but insufficient because researchers may privately prefer alternatives while publicly conforming to community incentives
Epistemic debt	Collective knowledge base	Hidden assumptions accumulate because dominant practices go unchallenged	Long-term reliance on shared representations, metrics, or benchmarks	Capacity to question background assumptions	Clarifies one downstream consequence of lock-in: assumptions become invisible because everyone shares them
Scaffolding dependence	Research infrastructure ecosystem	Support systems built for efficiency later	Standardized datasets, toolkits, benchmark platforms	Reversibility of methodological choices	Shows how infrastructure designed to

		constrain change			accelerate research can deepen lock-in by making deviation costly
--	--	---------------------	--	--	--

These distinctions clarify why scientific lock-in deserves recognition as a distinct failure mode. It is not simply “history matters” (path dependence), nor “the software is sticky” (technical lock-in), nor “I’m used to it” (cognitive lock-in). It is the process by which early AI adoption in materials domains creates a self-perpetuating monopoly on attention, resources, and legitimacy that subsequent innovation must overcome. Understanding this failure mode is therefore foundational to the remainder of the analysis.

Mechanisms of Lock-In

Four interlocking mechanisms convert early AI adoption into durable scientific lock-in within materials communities.

Increasing returns

Early-adopted methods rapidly attract disproportionate resources—citations, tool development, dataset curation, and talent—creating a compounding advantage that later entrants cannot match [1, 15]. In materials AI, the early success of graph-network architectures for crystal property prediction [5] generated a cascade of follow-on work: open-source implementations proliferated, benchmark datasets were curated around the same architecture, and review articles canonized it as the default. Each new publication citing the original work further entrenches the method, producing the classic increasing-returns dynamic in which “success breeds success.”

Switching costs

Once a pipeline is established—data featurized according to one scheme, models trained on one architecture, results benchmarked against one metric—replicating the entire workflow with an alternative incurs prohibitive costs in time, computational resources, and human expertise [21, 22]. Materials researchers who have invested years in curating datasets compatible with a particular descriptor set face substantial re-engineering costs to adopt a new representation, even if the new representation offers superior generalizability. These switching costs are not merely financial; they include the opportunity cost of delayed publications and the reputational risk of appearing to abandon a previously successful line of work.

Network effects

The value of an AI method increases with the number of researchers using it. When the majority of the community converges on a single framework, coordination benefits—shared code repositories, standardized APIs, joint workshops, and common review criteria—accrue to participants. At the same time, non-participants incur fragmentation costs [20]. In materials AI, network effects are visible in the rapid standardization around certain graph-network libraries; researchers adopting alternatives must either translate their results into the dominant format or risk being ignored in collective discussions and meta-analyses.

Institutionalization

Early methods become embedded in the social infrastructure of the field: graduate curricula teach them as standard, funding calls reference them explicitly, hiring committees favor candidates with demonstrated expertise in them, and high-impact journals develop implicit preferences for results produced with them [13, 18]. Institutionalization transforms a contingent early choice into an expected norm, so that deviating from the dominant paradigm requires explicit justification while conforming to it does not.

Conceptually, these mechanisms form a self-reinforcing cycle: early adoption triggers increasing returns that lower the relative cost of further adoption (Mechanism 1), which raises switching costs for alternatives (Mechanism 2), which strengthens network effects around the dominant method (Mechanism 3), which in turn accelerates institutionalization (Mechanism 4), completing the loop and making escape progressively harder. The cycle is particularly potent in materials AI because the high cost of experimental validation makes computational communities especially reliant on shared methodological scaffolding. Once the scaffolding is in place, the community optimizes within it rather than questioning its foundations.

Types of Scientific Lock-In

Scientific lock-in in materials AI manifests in four distinct but interrelated types.

Representational lock-in

This occurs when an early choice of descriptors or a featurization scheme becomes the universal language through which all subsequent modeling is conducted. Once the community converges on a particular set of structural or electronic descriptors, models are built, benchmarks are defined, and physical interpretations are derived exclusively within that representational space [3, 5]. The mechanism is primarily increasing returns: the initial descriptor set accrues validation data, error analyses, and interpretability tools that no competing representation can match. A materials example is the widespread early adoption of certain graph-network node and edge features for crystals; alternative featurizations, even those shown to capture long-range interactions more efficiently, face an uphill battle because the entire validation ecosystem is calibrated to the original choice.

Methodological lock-in

Here, the community coalesces around a particular algorithm family—e.g., graph neural networks versus kernel methods or random forests—such that the dominant family defines the intellectual boundaries of the subfield [2]. The mechanism combines network effects and institutionalization: conferences, special issues, and funding priorities become organized around the dominant method, creating a de facto monopoly on legitimacy. In solid-state materials science, the early embrace of graph networks [5] has made alternative methodological families appear niche even when they offer complementary strengths in data-scarce regimes.

Data lock-in

Early selection of a particular database or dataset as the community standard creates lock-in because all benchmarking, transfer learning, and model comparison become tethered to that dataset's biases and limitations [3]. The mechanism is switching cost plus network effects: rebuilding benchmarks on a new dataset requires coordinated community effort that no single laboratory can undertake alone. Materials examples include the canonical status of certain high-throughput computational databases in property-prediction tasks; even when newer, more diverse datasets emerge, researchers continue to report results on the legacy set because comparability with prior literature is prized.

Evaluation lock-in

The community adopts an early benchmark metric or validation protocol that subsequently defines “success” for the entire field [4]. The mechanism is institutionalization: peer review and funding decisions implicitly reward models that excel on the entrenched metric, even if that metric is known to be incomplete or misleading for real-world materials performance. In inverse-design tasks, for instance, early reliance on certain formation-energy or band-gap metrics has shaped expectations such that methods excelling on alternative figures of merit (e.g., synthesizability or stability under operating conditions) are undervalued.

Each type reinforces the others. Representational lock-in feeds methodological lock-in by limiting the hypothesis space that algorithms can explore; data lock-in cements evaluation lock-in by providing the ground truth against which all models are judged. Together they create a coherent epistemic cage whose bars are invisible precisely because they are shared by the entire community.

A Typology of Lock-In Failure Modes

Lock-in does not merely constrain options; it produces four characteristic failure modes that undermine the scientific integrity of materials AI. An inferior method, representation, or dataset continues to dominate despite the availability of clearly superior alternatives. The mechanism is the full lock-in cycle: increasing returns and institutionalization outweigh the recognized disadvantages. In materials AI, a featurization scheme known to miss long-range interactions may persist because the entire citation and benchmarking ecosystem is built around it [5]. Detection signature: persistent use of the method in new papers even after comparative studies demonstrate superior performance elsewhere.

Promising new AI approaches fail to gain traction, not because they lack merit but because they fall outside the locked-in paradigm. The mechanism is network effects combined with high switching costs; reviewers and editors demand comparability with the dominant framework, effectively requiring innovators to solve an additional translation problem before their core contribution can be assessed [20]. A materials example is the delayed uptake of alternative architectures that excel in sparse-data regimes because the community's evaluation infrastructure is optimized for dense graph-network benchmarks. Detection signature: repeated rejection or marginalization of papers introducing non-dominant methods despite strong technical validation. The community loses the ability to perform fair, head-to-head comparisons across methodological families because the dominant approach has become the only one with standardized infrastructure. The mechanism is evaluation lock-in: without common ground, new methods are either forced into the dominant mold (losing their unique strengths) or dismissed as incomparable. In practice, this means materials researchers cannot reliably determine whether graph networks truly outperform kernel methods across the full materials space [3]. Detection signature: systematic absence of ablation studies or cross-family benchmarks in the literature. The community collectively recognizes that lock-in has occurred and that better alternatives exist, yet remains unable to coordinate a switch. The mechanism is institutionalization at the highest level: funding agencies, journals, and training programs have all internalized the dominant paradigm, creating a coordination dilemma akin to a multi-player prisoner's dilemma [11, 16]. Materials AI exhibits early warning signs of collective regret when review articles begin to note the limitations of canonical methods while simultaneously acknowledging their continued dominance. Detection signature: proliferation of opinion pieces and meta-discussions lamenting the status quo without accompanying changes in practice.

These failure modes are not hypothetical; they are the logical endpoints of the mechanisms and types articulated earlier. By naming and classifying them, the field gains a diagnostic vocabulary with which to recognize lock-in before it becomes irreversible.

Detection Principles

Detecting scientific lock-in before it becomes irreversible requires systematic attention to observable community-level indicators rather than isolated modeling outcomes. Four detection principles provide a practical framework for materials AI researchers and institutions to identify when early adoption has begun to constrain epistemic possibilities.

Method concentration

The first and most direct indicator is an unusually high concentration of publications, tools, and datasets around a single AI method or representation within a given subfield. When the Herfindahl-Hirschman Index applied to methodological choices exceeds typical thresholds observed in other scientific domains, it signals that increasing returns have begun to dominate [3, 5]. In materials AI, for instance, one can track the proportion of papers in npj Computational Materials or

Chemistry of Materials that rely on the same graph-network architecture for crystal modeling; persistent dominance above 70%–80% over multiple years suggests representational or methodological lock-in rather than organic convergence on a genuinely optimal solution. This principle is powerful because it is quantifiable through bibliometric analysis and can be monitored longitudinally without requiring subjective judgment.

Citation bias

Early-adopted papers receive citations at rates that cannot be explained by their intrinsic technical merit alone. When foundational works on a particular AI framework continue to dominate reference lists even after newer studies demonstrate clear shortcomings, citation bias reveals the self-reinforcing loop of increasing returns and network effects [2, 15]. Materials researchers can detect this by examining citation half-lives. If the seminal graph-network paper from 2019 [5] is still cited more frequently than all subsequent methodological innovations combined, the community is likely experiencing lock-in. The bias is especially telling when it persists across review articles and funding proposals, where authors cite the dominant method performatively to signal membership in the community rather than to engage its substance.

Switching resistance

Attempts to introduce alternative AI approaches encounter disproportionate skepticism or rejection during peer review, conference presentations, or grant evaluation. Switching resistance manifests as demands that new methods be benchmarked exclusively against the locked-in standard, effectively forcing innovators to solve the dominant paradigm's problems before their own contributions can be assessed [20, 21]. In materials AI, this appears when proposals exploring alternative featurization schemes or non-graph architectures are critiqued not on their merits but on their lack of “compatibility” with established pipelines. Detection involves cataloging reviewer comments and acceptance rates for non-dominant submissions; systematic patterns of resistance indicate that network effects and institutionalization have raised the bar for alternatives beyond what technical merit alone would justify.

Comparative gap

The literature exhibits a systematic absence of rigorous head-to-head comparisons between the dominant AI approach and plausible alternatives. When papers report performance only within the locked-in framework and omit ablation studies or cross-family benchmarks, the community has lost the ability to assess true relative value [3, 4]. In solid-state materials science, this gap is evident in the scarcity of studies that test graph networks against kernel methods or transformer architectures on identical datasets; the absence is not accidental but a symptom of evaluation lock-in, where the dominant metric and dataset have become the only accepted ground truth. Tracking this principle through meta-analyses of recent publications reveals whether the field is operating in a state of comparative ignorance.

Taken together, these four principles function as an early-warning system. Regular application—through community-maintained dashboards or periodic meta-reviews—can flag lock-in while escape is still feasible, before institutionalization makes reversal prohibitively expensive. Because scientific lock-in is a collective phenomenon, detection must itself be collective; individual researchers applying these principles in isolation will still face the very coordination costs the principles aim to diagnose [1, 16].

Mitigation Principles

Preventing or escaping scientific lock-in demands proactive, community-wide interventions that deliberately counteract the four mechanisms identified earlier. Five mitigation principles translate the conceptual analysis into actionable strategies for materials AI.

Research groups, funding agencies, and journals should deliberately maintain parallel streams of inquiry using multiple AI representations and algorithmic families rather than converging on a single dominant approach [2, 3]. This principle directly counters increasing returns by ensuring that no single method accumulates an unassailable resource advantage. In practice, pluralism can be implemented through dedicated tracks in conferences for “alternative paradigms” or by requiring grant proposals to justify why only one methodological family is pursued when complementary approaches exist. Over time, pluralism keeps switching costs low because multiple toolchains remain actively maintained and familiar to the community. Every new materials AI contribution must include explicit, standardized comparisons against at least two non-dominant alternatives using identical datasets and evaluation protocols [4, 5]. This principle attacks evaluation lock-in at its root by making comparative ignorance impossible. Journals such as *Machine Learning: Science and Technology* could enforce this through checklist requirements. At the same time, benchmark platforms could host modular leaderboards that display performance across representational families rather than within a single locked-in metric. Standardized comparisons force the community to confront trade-offs openly rather than allowing the dominant method to define success by default. Infrastructure developers should prioritize modular, interoperable systems that allow researchers to swap representations, algorithms, or datasets with minimal re-engineering [21, 22]. Open-source libraries that expose common interfaces for featurization and model evaluation exemplify this principle; a researcher locked into one graph-network implementation could then test an alternative descriptor set by changing only a configuration file rather than rebuilding the entire pipeline. By design, such modularity reduces the technical and cognitive barriers that amplify network effects and institutionalization.

The materials AI community should institutionalize periodic, independent audits of methodological diversity using the detection principles outlined above [11, 16]. Every three to five years, a cross-institutional working group could publish a “state of lock-in” report that quantifies method concentration, citation bias, and comparative gaps across major subfields. These audits create transparency and generate the shared knowledge necessary to overcome coordination dilemmas inherent in collective regret. Funding bodies must allocate a non-trivial fraction of resources—perhaps 20%–30%—explicitly to the development and benchmarking of non-dominant AI methods, representations, and datasets [13, 18]. Targeted calls for “methodological challengers” prevent first-mover advantages from becoming permanent by ensuring that alternatives receive the resources needed to reach maturity. This principle is especially potent because it interrupts the institutionalization mechanism at the source: when hiring, promotion, and tenure decisions begin to value demonstrated work on alternative paradigms, the community’s incentive structure realigns toward pluralism.

Table 2 consolidates the paper’s core framework by linking each lock-in type to its principal object of entrenchment, its most visible failure signature, and the intervention lever most likely to preserve long-term adaptability.

Table 2. Cross-level architecture of scientific lock-in in materials AI: from entrenchment object to failure signature and intervention lever

Lock-in type	Primary object of entrenchment	Dominant reinforcing mechanisms	Characteristic epistemic consequence	Most visible detection signature	Highest-leverage mitigation lever
Representational lock-in	Descriptor set, featurization logic, representational language	Increasing returns; switching costs	The field can only “see” materials problems through one modeling vocabulary, narrowing hypothesis space	Persistent reuse of the same descriptors even after alternatives show superior coverage or interpretability	Maintain parallel representational pipelines and modular featurization interfaces

Methodological lock-in	Algorithm family or modeling paradigm	Network effects; institutionalization	One method family defines legitimacy, making alternatives appear niche before they are evaluated on equal terms	Publication, workshop, and review dominance by one model family	Require cross-family benchmarking and reviewer prompts that assess alternatives explicitly
Data lock-in	Benchmark dataset or canonical database	Switching costs; network effects	Community inference becomes tethered to the biases and omissions of legacy datasets	Continued reliance on legacy benchmarks despite the emergence of broader or more realistic datasets	Fund challenger datasets and require dual-benchmark reporting
Evaluation lock-in	Metric, validation protocol, and leaderboard logic	Institutionalization	Success becomes defined by what the entrenched metric can capture, not by what matters scientifically or practically	Sparse use of alternative metrics and absence of benchmark designs aligned with real materials performance	Expand reporting standards to include multi-metric, application-relevant evaluation
System-level result	The entire materials AI ecosystem	Joint action of all four mechanisms	Comparative openness declines, and the field risks suboptimal persistence, innovation suppression, comparative ignorance, and collective regret	High method concentration, citation bias, switching resistance, and comparative gaps appear together	Combine pluralism, standardization, modularity, audits, and targeted funding rather than relying on single-point fixes

Collectively, these mitigation principles do not seek to eliminate early adoption—itself a necessary driver of progress—but to ensure that early choices remain revisable. By embedding pluralism, comparability, modularity, transparency, and balanced resourcing into the fabric of materials AI, the field can enjoy the productivity gains of rapid innovation while preserving the long-term adaptability that defines scientific health [1, 15].

Relation to Other Failure Modes

Scientific lock-in does not exist in isolation; it interacts with, amplifies, and is amplified by other recognized failure modes in technology-intensive science.

First, lock-in is the concrete realization of path dependence at the community scale. While path dependence describes any historical contingency that shapes future trajectories [11, 12], scientific lock-in is the pathological subset in which that contingency produces self-reinforcing sub-optimality. The lock-in cycle can be conceptualized as follows: an initial early

adoption of a specific AI method or representation triggers increasing returns that compound usage and visibility; these returns raise switching costs and strengthen network effects; the resulting coordination advantages accelerate institutionalization through curricula, funding, and hiring; and institutionalization in turn makes future adoption decisions even more heavily weighted toward the original choice, closing the loop and converting contingent history into durable constraint. This cycle illustrates how path dependence becomes lock-in when social and epistemic feedback mechanisms are allowed to operate unchecked.

Second, scientific lock-in generates epistemic debt. Once a community commits to a locked-in representation or evaluation metric, underlying assumptions—such as the sufficiency of local atomic environments or the adequacy of a particular formation-energy benchmark—become invisible because they are shared by every participant [3, 4]. New researchers inherit these assumptions without ever having examined them, accumulating a form of collective epistemic debt that only becomes apparent when an external shock (such as experimental failure of a predicted material) forces re-examination. Lock-in hides the debt by removing the very comparisons that would reveal it.

Third, lock-in is reinforced by the scaffolding problem, in which the very infrastructure built to accelerate research—standardized datasets, open-source toolkits, benchmark platforms—becomes the cage that prevents escape [16, 20]. In materials AI, the scaffolding of canonical high-throughput databases and graph-network libraries was erected to solve immediate data-scarcity challenges. Yet, that same scaffolding now enforces data lock-in and methodological lock-in by making deviation computationally and socially expensive. The relation is mutually constitutive: lock-in creates the demand for more scaffolding, and the scaffolding deepens lock-in.

Recognizing these interrelations equips the community to address lock-in not as a standalone technical issue but as part of a broader ecosystem of failure modes. Mitigation strategies that target only lock-in in isolation will be undermined unless they simultaneously reduce path dependence, retire epistemic debt, and redesign scaffolding for reversibility. The framework developed here, therefore, serves as a diagnostic lens through which other failure modes in materials AI can be more clearly understood and more effectively managed [1].

Implications for Materials AI Practice

The identification of scientific lock-in carries direct, actionable implications for three stakeholder groups whose practices collectively shape the field's trajectory.

For individual authors, three changes are essential. First, every manuscript must include explicit comparisons against at least one non-dominant alternative, even when space is limited; this habit prevents the inadvertent reinforcement of evaluation lock-in [5]. Second, researchers should avoid method monomania by presenting results from multiple representational families within the same study, thereby demonstrating robustness rather than merely reporting performance within the locked-in paradigm [2]. Third, authors must explicitly acknowledge lock-in risks in the discussion section—citing the possibility that observed success may reflect community investment rather than intrinsic superiority—thereby normalizing reflexive awareness [1].

For reviewers and editors, the implications center on shifting evaluative criteria. Reviewers should routinely ask whether a submission has engaged with plausible alternatives and whether the chosen metric or dataset is itself locked in; manuscripts that fail to address these questions should be returned for revision rather than rejected outright, creating space for methodological pluralism [3, 4]. Editors can further institutionalize this expectation by updating author guidelines and reviewer checklists, transforming peer review from a gatekeeper of conformity into a guardian of epistemic diversity.

For the broader community—funding agencies, conference organizers, and curriculum designers—the implications demand structural intervention. Funding calls should reserve portions for alternative-paradigm research; conferences should host dedicated sessions on non-dominant methods; and graduate curricula should teach at least two competing AI frameworks side-by-side rather than presenting a single canonical pipeline as the state of the art [13, 18]. These

community-level changes reduce switching costs, weaken network effects, and prevent institutionalization from ossifying around any single early choice.

Taken together, these practice-level shifts transform scientific lock-in from an invisible background risk into a manageable design parameter. Materials AI can retain the enormous productivity benefits of rapid early adoption while safeguarding the field's long-term capacity for self-correction and genuine discovery [15, 27-29]. The ultimate implication is that methodological pluralism is not a luxury but a necessary condition for sustainable progress in a data-driven, computationally intensive science.

Conclusion

Scientific lock-in through early AI adoption in materials domains constitutes a distinct and previously under-examined failure mode whose consequences extend far beyond any single suboptimal model or dataset. By defining the phenomenon, dissecting its mechanisms, classifying its types, and mapping its resultant failure modes, this analysis has shown how an initial burst of productive innovation can inadvertently harden into epistemic rigidity that suppresses future progress. The self-reinforcing cycle—early adoption feeding increasing returns, raising switching costs, strengthening network effects, and culminating in institutionalization—threatens to trade short-term gains in predictive accuracy for long-term losses in scientific adaptability.

The framework offered here, therefore, calls for deliberate lock-in awareness across the materials AI community. Researchers, reviewers, funders, and educators must treat early adoption decisions not as neutral technical choices but as high-stakes commitments whose downstream consequences require explicit management. Methodological pluralism, standardized comparisons, modular infrastructure, regular audits, and balanced funding for alternatives together provide a practical pathway to preserve the field's capacity for self-correction.

Only by maintaining deliberate openness to alternative representations, algorithms, and evaluation paradigms can materials AI fulfill its promise as a genuinely transformative scientific enterprise rather than a self-limiting echo chamber of early success. The time to act is now, while the field's rapid growth still affords the flexibility to choose pluralism over path dependence. The future of materials discovery depends on our willingness to keep the epistemic doors open even as we accelerate through them.

Acknowledgements

None

Conflict of interest

None

Financial support

None

Ethics statement

None

Published online: 18 July 2025

Rights and permissions

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

References

- Qiu TA, He Z, Chugh T, Kleiman-Weiner M. The lock-in hypothesis: Stagnation by algorithm. arXiv preprint arXiv:2506.06166. 2025 Jun 6.
- 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.
- Zunger A. Inverse design in search of materials with target functionalities. *Nat Rev Chem*. 2018;2(4):0121.
- Chen C, Ye W, Zuo Y, Zheng C, Ong SP. Graph networks as a universal machine learning framework for molecules and crystals. *Chem Mater*. 2019;31(9):3564-72.
- Vasudevan RK, Choudhary K, Mehta A, Smith R, Kusne G, Tavazza F, et al. Materials science in the artificial intelligence age: High-throughput library generation, machine learning, and a pathway from correlations to the underpinning physics. *MRS Commun*. 2019;9(3):821-38.
- Goodge BH, El Baggari I, Hong SS, Wang Z, Schlom DG, Hwang HY, et al. Disentangling coexisting structural order through phase lock-in analysis of atomic-resolution STEM data. *Microsc Microanal*. 2022;28(2):404-11.
- Hansen M. From attention economy to cognitive lock-ins. *Big Data Soc*. 2024;11(3):20539517241275878.
- Gerber S, Yang SL, Zhu D, Soifer H, Sobota JA, Rebec S, et al. Femtosecond electron-phonon lock-in by photoemission and x-ray free-electron laser. *Science*. 2017;357(6346):71-5.
- Chiodo M, Müller D, Sienknecht M. Educating AI developers to prevent harmful path dependency in AI resort-to-force decision making. *Aust J Int Aff*. 2024;78(2):210-9.
- Chu H, Hassink R, Martin R, Sunley P, Unruh G. Rethinking path dependence and lock-ins in regions, economy and society. *Camb J Reg Econ Soc*. 2026:rsag001.
- Camuffo A, Gambardella A, Kazemi S. A theory based analysis of path dependence. *Econ Innov New Technol*. 2025:1-24.
- Sauber MH. Technology upgrading or path dependence? Malaysia's AI infrastructure strategy in regional perspective. *Asian J Technol Innov*. 2025:1-35.

Raineau L. Rethinking path dependence, technical innovation and social practices in a renewable energy future. *Energy Res Soc Sci.* 2022;84:102411.

Nafizah UY, Roper S, Mole K. Estimating the innovation benefits of first-mover and second-mover strategies when micro-businesses adopt artificial intelligence and machine learning. *Small Bus Econ.* 2024;62(1):411-34.

Helmrich A, Chester M, Miller TR, Allenby B. Lock-in: Origination and significance within infrastructure systems. *Environ Res Infrastruct Sustain.* 2023;3(3):032001.

Zhang Z, Yu J, Hu C. Evolution of path dependence and lock-in with the knowledge mining analysis of web of science literature. *J Intell Fuzzy Syst.* 2018;35(3):2951-63.

Xu Y, Wu B, Shuai L. Is digital technology the panacea for alleviating carbon lock-in in manufacturing? *Environ Dev Sustain.* 2025;1-32.

Chen P, David KG, Zhao D. Identifying technology lock-in and tracing knowledge source trajectories: A case study of lithography. *Technol Anal Strateg Manag.* 2024;36(10):2414-29.

Seidel VP, Langner B, Sims J. Dominant communities and dominant designs: Community-based innovation in the context of the technology life cycle. *Strateg Organ.* 2017;15(2):220-41.

Đerić E, Frank D, Tomić Z. Factors influencing perceived switching cost and user switching behavior towards AI-based solutions and technologies-a systematic review. In: 2023 IEEE 21st Jubilee International Symposium on Intelligent Systems and Informatics. SISY. Piscataway, NJ: IEEE; 2023. p. 415-20.

Zhou T, Li S. Examining user switching intention between generative AI platforms: A push-pull-mooring perspective. *Inf Dev.* 2025;41(3):692-704.

Liu Y, Ji P. Understanding designers' switching intention to AI painting tools using the PPM framework. *Sci Rep.* 2025;15(1):2246.

Ashby NJ, Teodorescu K. The effect of switching costs on choice-inertia and its consequences. *PLoS One.* 2019;14(3):e0214098.

Wang J, Yu Y, Lu C. Switching intention of potential smart home users: The co-action of innovation diffusion theory and switching costs. *Int J Hum Comput Interact.* 2025;41(5):3509-26.

Huichao R, Feng D. Digital technology lock-in: A study of erosion effects and unlocking strategies from the reverse locking perspective. *Sci Technol Prog Policy.* 2025;42(2):1-10.

Yang S, Sun Y. The first-mover advantage in FinTech: Evidence from the banking industry in China. *Appl Econ Lett.* 2025;1-5.

Lee D, Kim Y. First-mover advantages and intentional knowledge spillover effects on cybersecurity start-ups' financial and innovation performance: Incentive for revealing innovations in high-tech emerging industry. *Asian J Technol Innov.* 2024;32(1):1-9.

Bolender B, Vispoel S, Converse G, Koprowicz N, Song D, Osaro S. Generative AI in K12: Analytics from early adoption. *J Meas Eval Educ Psych.* 2024;15(Spec Issue):361-77.