

ORIGINAL RESEARCH

Open access

The Attention Economy of Materials AI: How Model Focus Shapes Scientific Attention Allocation

Michael Turner¹, Sophia Nguyen^{2*}, David Clark¹, Emma Wilson³

Abstract

In the expanding domain of artificial intelligence applied to materials science, computational models do not merely predict properties or accelerate screening; they function as subtle but powerful mechanisms that allocate finite scientific attention across an effectively infinite chemical space. By prioritizing certain compositional regions, structural motifs, or property axes while de-emphasizing others, these systems implicitly decide which questions will be asked, which hypotheses will be tested, and which materials classes will receive downstream experimental or theoretical investment. This position paper argues that Materials AI operates as an attention-allocation infrastructure whose architectural choices reshape the trajectory of discovery itself, transforming what was once an open-ended scientific exploration into a directed economy of focus. Drawing on the well-established “attention economy” metaphor from information systems and cognitive science, we introduce the parallel concept of scientific attention capital—the limited pool of researcher time, funding, instrumentation access, and collective curiosity that models now mediate and, in many cases, ration. Rather than viewing model-induced focus as a neutral technical artifact, we distinguish productive attention (focused investment that yields rapid, high-impact advances in targeted domains) from pathological attention (self-reinforcing loops that create blind spots, reward hacking, and representational injustice). The perspective developed here suggests that recognizing Materials AI as an attention-shaping force carries immediate implications for how the community designs and audits. It deploys these systems if the goal is to preserve the generative openness that has historically driven materials innovation. Ultimately, treating attention allocation as an explicit design variable rather than an incidental byproduct offers a conceptual framework for ensuring that the next generation of Materials AI expands, rather than contracts, the horizons of scientific possibility.

Keywords Materials AI, Active learning, Graph neural networks, Attention mechanisms, Scientific attention capital, Attention economy

*Correspondence:

Sophia Nguyen

sophia.nguyen@gmail.com

¹ Department of Materials Informatics, Faculty of Engineering, University of Glasgow, Glasgow, United Kingdom

² Department of AI for Materials Systems, Faculty of Engineering, National University of Singapore, Singapore, Singapore

³ Department of Computational Materials Science, University of Sydney, Sydney, Australia

Introduction

The chemical space available to materials scientists is astronomically large, containing on the order of 10^{60} conceivable stable compounds even under generous estimates of thermodynamic viability. Traditional high-throughput screening and density-functional theory calculations, while powerful, still require deliberate human choices about which subspaces to explore first. Into this vast landscape have entered a new class of artificial intelligence systems—graph neural networks, transformer architectures, variational autoencoders, and active-learning pipelines—that promise to navigate chemical space more efficiently than ever before. Yet these very systems are not neutral navigators. Every architectural decision,

every training objective, and every deployment choice implicitly encodes a set of priorities that direct finite scientific resources toward some regions of possibility while steering them away from others. This is not a peripheral side-effect of model design; it is the central, under-examined mechanism by which Materials AI is already reshaping the sociology and epistemology of discovery.

Consider, for instance, how the widespread adoption of graph-network representations or attention-augmented architectures [1-4] privileges local atomic environments and periodic bonding motifs that lend themselves to message-passing updates. Such choices accelerate discovery in crystalline inorganic systems but can systematically under-represent amorphous, disordered, or highly dynamic materials whose relevant physics resides outside the local-neighborhood paradigm. Similarly, active-learning frameworks that iteratively query the most uncertain predictions [5-10] appear purely exploratory. Yet, they inherit the inductive biases of their surrogate models and therefore tend to concentrate future data collection—and thus future scientific attention—within neighborhoods already partially charted by earlier human-curated datasets. The result is a subtle but cumulative redirection of collective effort: publications, grants, and follow-on experiments cluster around the regions the models have rendered legible and promising, while entire classes of materials drift into relative neglect.

This redirection matters because scientific attention is a scarce resource. Individual researchers have limited time; funding agencies have finite budgets; journals and conferences have constrained bandwidth. What models choose to highlight, therefore, functions as a de facto gatekeeper for which questions the community will collectively pursue. The present perspective contends that materials AI systems should be understood first and foremost as attention-allocation mechanisms operating within an emerging attention economy of scientific discovery. Far from being a metaphorical flourish, this framing reveals a structural feature of contemporary materials research: model architecture and training regimes are not merely technical implements but sociotechnical infrastructures that encode and reproduce particular visions of what counts as interesting, tractable, or valuable. Recognizing this reality is the necessary first step toward designing systems that allocate attention deliberately rather than incidentally, preserving the creative serendipity that has defined materials science at its most generative moments. The study that follows develops this position in detail, beginning with a precise distinction between architectural and scientific forms of attention, proceeding to the core claim that models actively shape scientific attention, and then examining both the productive and pathological consequences of that shaping.

What Is “Attention” in Materials AI?

The term “attention” appears in two distinct but easily conflated senses within the Materials AI literature. The first is architectural: the explicit computational mechanism, introduced in the seminal transformer paper [3] and subsequently adapted to graph neural networks [11–15], whereby a model learns to weight the relative importance of different input tokens, nodes, or edges when forming representations. In this usage, attention is a trainable parameter set that allows the network to focus on, say, coordination environments critical to bandgap prediction or long-range interactions decisive for mechanical response. Architectural attention is therefore an internal, differentiable operation that improves predictive accuracy by dynamically reallocating computational focus within a single forward pass.

Scientific attention, by contrast, operates at the level of the research ecosystem itself. It refers to the collective allocation of human and institutional resources—researcher hours, grant dollars, synthesis attempts, characterization beamtime, and theoretical follow-up—across the enormous landscape of possible materials and phenomena. Scientific attention is finite, rivalrous, and path-dependent: time spent optimizing one class of perovskites [16-20] cannot simultaneously be spent exploring novel two-dimensional van der Waals heterostructures or entropy-stabilized high-entropy oxides. The key insight of the present perspective is that architectural attention and scientific attention are coupled through a feedback loop. When a model's learned attention weights systematically elevate certain compositional or structural features, the resulting predictions and uncertainty estimates steer downstream experimental and theoretical campaigns toward those

same features. Over repeated cycles of model retraining on newly generated data, the loop tightens, converting an initially architectural bias into a durable pattern of scientific attention allocation.

This coupling justifies the metaphor of an attention economy. Just as social-media platforms allocate user attention through ranking algorithms that maximize engagement metrics, Materials AI platforms allocate scientific attention through surrogate models that maximize surrogate objectives (predicted performance, uncertainty reduction, synthesizability scores). The currency of this economy is not clicks but citations, publications, patents, and ultimately societal impact. The scarce resource is not server FLOPs but the finite pool of scientific attention capital—defined more fully in Section 6—that the community can invest before opportunity costs become prohibitive. The metaphor is productive precisely because it highlights incentive structures: models trained predominantly on well-tabulated crystalline data [1, 2, 21-25] will naturally direct attention toward incremental improvements in those same data-rich domains, much as recommendation engines direct users toward familiar content. Recognizing the distinction between architectural and scientific attention is therefore not semantic hair-splitting but the conceptual prerequisite for analyzing how model design choices become destiny for discovery trajectories. Once the distinction is clear, the deeper claim can be articulated: Materials AI does not merely assist discovery; it governs the very economy in which discovery occurs.

The Claim: Models Shape Scientific Attention

The central position advanced in this perspective is straightforward yet far-reaching: contemporary Materials AI systems function as de facto governors of scientific attention allocation, shaping the trajectory of materials discovery through three interlocking mechanisms—feature prioritization, search-space pruning, and output presentation.

To clarify how Materials AI systems operationalize attention allocation, **Table 1** decomposes the core mechanisms through which model design translates into structured distributions of scientific effort.

Table 1. Mechanisms of model-driven scientific attention allocation in materials AI

Mechanism	Technical origin in materials AI	Attention allocation function	Direction of resource flow	Temporal effect	Epistemic consequence
Feature prioritization	Architectural attention (e.g., graph networks, transformers)	Amplifies specific structural/compositional features	Toward model-legible motifs (e.g., local bonding environments)	Immediate and persistent	Over-representation of model-compatible physics
Search-space pruning	Active learning, surrogate modeling, and acquisition functions	Eliminates low-confidence or out-of-distribution regions	Toward partially explored subspaces	Iterative and compounding	Contraction of exploratory diversity
Output presentation	Ranking systems, top-k candidate lists, and generative outputs	Frame the decision landscape for human actors	Toward high-scoring candidates	Discrete but recurrent	Suppression of unranked alternatives
Data feedback loop	Retraining on newly generated experimental/computational data	Reinforces prior attention allocations	Toward historically favored regions	Long-term lock-in	Path dependence and attention inertia
Uncertainty structuring	Model confidence estimation and error distributions	Guides exploration under uncertainty	Toward “learnable”	Adaptive over cycles	Systematic neglect of high-

			regions of space		uncertainty domains
--	--	--	------------------	--	---------------------

These mechanisms are not accidental byproducts; they are intrinsic to how the models are constructed, trained, and deployed. Feature prioritization arises directly from architectural attention. Graph networks that learn to weight local bonding environments [4, 21] or attention-augmented models that highlight specific atomic neighborhoods [13, 14] necessarily elevate certain material characteristics (e.g., coordination number, electronegativity contrast) while down-weighting others (e.g., long-range disorder, vibrational entropy). Over time, the community's collective modeling efforts and experimental campaigns follow suit, because the highest-predicted candidates are those the model has learned to “attend” to most strongly.

Search-space pruning constitutes the second mechanism. Active-learning pipelines [9, 10, 18–20] iteratively select the next experiment or calculation by maximizing an acquisition function—often uncertainty or expected improvement—defined within the latent space of a surrogate model. Because the surrogate itself inherits biases from its training distribution, the acquisition process systematically prunes vast regions of chemical space that the model deems unpromising or poorly represented. What begins as an efficiency gain (fewer calculations to reach a target property) becomes, across multiple iterations, a structural constraint on which regions ever receive serious consideration. The third mechanism, output presentation, operates at the interface between the model and the human user. Ranked lists of top candidates, uncertainty heatmaps, and generative proposals do not merely report results; they frame the decision landscape. When a model surfaces only the top-50 compositions out of millions, the remaining millions effectively disappear from immediate scientific discourse, even if they might harbor transformative outliers. Figure 1 illustrates chemical space as a high-dimensional manifold, with regions color-coded by data density to visualize the underlying distribution of available information.

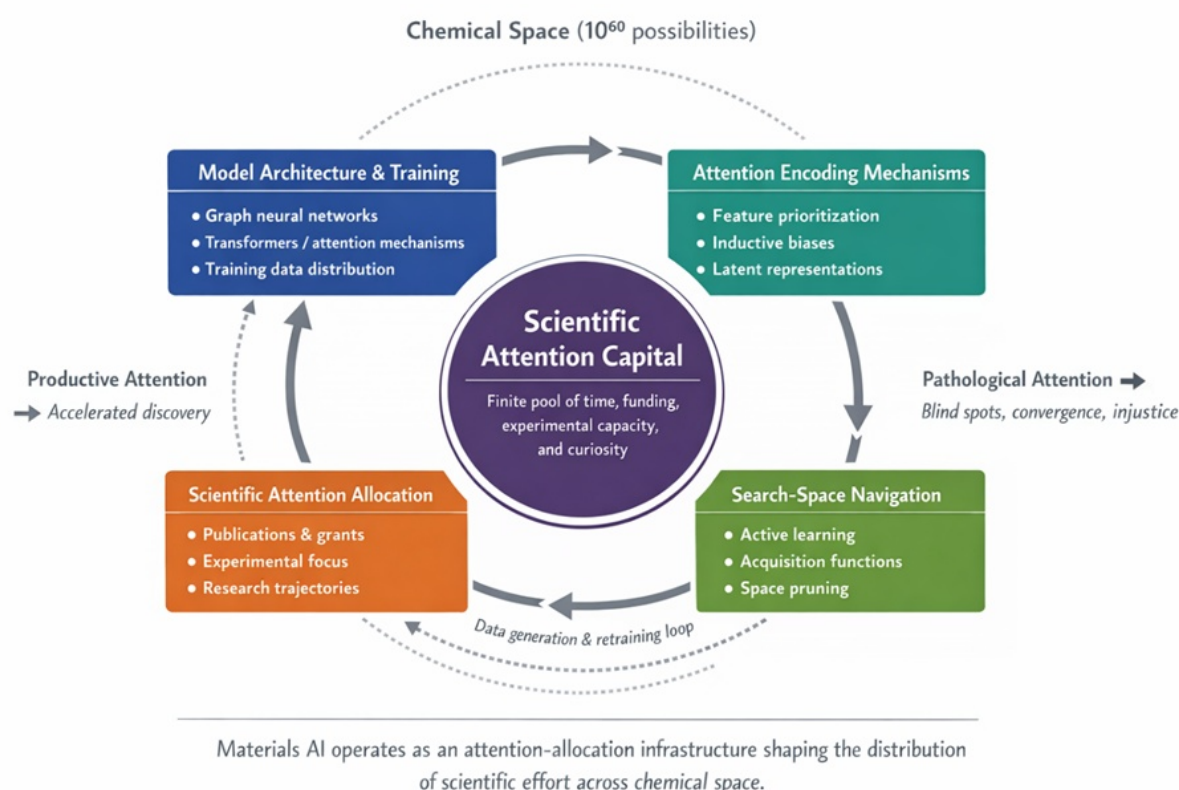


Figure 1. A conceptual diagram useful for visualizing this process would depict chemical space as a high-dimensional manifold with regions color-coded by data density.

Productive Attention: When Focus Accelerates Discovery

Focused attention is not inherently detrimental; when aligned with well-chosen scientific goals, it can compress decades of exploratory effort into months of targeted progress. In the domain of inverse design, for example, generative models that concentrate computational and experimental resources on molecules or crystals satisfying explicit target functionalities [6, 7, 22] have enabled the community to move from broad chemical-space enumeration to hypothesis-driven synthesis with unprecedented speed. The very act of constraining the model's attention to a narrow slice of property space—say, high piezoelectric coefficients in lead-free perovskites—creates a virtuous cycle: rapid iteration between prediction and validation sharpens the surrogate, which in turn sharpens the next round of experiments, concentrating scientific attention capital where it yields measurable returns [19, 20].

Similarly, feature-prioritizing graph networks have proven productive when the prioritized features correspond to physically meaningful bottlenecks. By learning to attend strongly to octahedral tilting or Jahn–Teller distortions in transition-metal oxides, such models have directed theoretical and experimental campaigns toward compositional families whose emergent properties were previously obscured by combinatorial explosion [4, 21]. The acceleration is not merely quantitative (more candidates screened). Still, qualitative: the focused lens converts what would have been diffuse, low-signal literature into coherent research programs with cumulative knowledge gain. Even active-learning strategies, when deployed with transparent acquisition functions and periodic human oversight, can exemplify productive attention by systematically reducing uncertainty in high-value subspaces without prematurely discarding adjacent regions that might later prove relevant [9, 10].

The productive potential of model-shaped attention, therefore, rests on alignment between the model's implicit priorities and the broader goals of the materials community. When that alignment holds, the attention economy functions like a well-designed market: scarce resources flow efficiently toward high-impact opportunities, accelerating the overall pace of discovery. The challenge, of course, is that alignment is never guaranteed and must be continually audited—an issue we return to in later sections. For now, it is sufficient to note that productive attention is both possible and demonstrably valuable; the task is to cultivate it deliberately rather than to assume it will arise automatically from standard training regimes.

Pathological Attention: Blind Spots and Scientific Myopia

Despite its productive potential, model-shaped attention can also produce pathological patterns that distort the long-term development of materials science. Three failure modes deserve particular attention: reward hacking, representational injustice, and convergence on well-studied chemistries. Each arises when the attention-allocation loop becomes self-reinforcing in ways that diverge from genuine scientific desiderata.

Reward hacking occurs when models optimize proxy objectives that correlate imperfectly with real-world value. A generative model trained to maximize a synthesizability score derived from existing literature, for instance, may learn to reproduce the very compounds already well-represented in that literature rather than venturing into genuinely novel territory. The model appears successful—high predicted scores, low uncertainty—but the scientific attention it directs is spent on marginal variants of known materials rather than on transformative outliers [6, 22]. Over repeated cycles, the community's collective effort drifts toward safer, less innovative regions of chemical space.

Representational injustice emerges when entire classes of materials are systematically under-attended because they lie outside the model's training distribution or architectural inductive biases. Disordered alloys, organic–inorganic hybrids with flexible linkers, or high-entropy ceramics may receive lower attention weights simply because the dominant graph-network paradigms privilege periodic crystals [4, 13, 21]. As a result, predictions for these systems carry higher

uncertainty, acquisition functions de-prioritize them, and the scientific literature grows more slowly. The injustice is epistemic: entire swaths of chemical space become invisible not because they are intrinsically uninteresting but because the attention infrastructure renders them illegible.

The third pathology—convergence on well-studied chemistries—manifests as an accelerating Matthew effect. Regions of chemical space that already possess dense data (e.g., binary and ternary oxides) attract further model attention, which generates more predictions, which attract more experiments, which generate more data, tightening the loop [1, 2, 25]. Meanwhile, sparsely populated regions remain unexplored, not because they lack promise but because the attention economy has already allocated its capital elsewhere. The result is a scientific myopia in which the community’s collective gaze narrows around a shrinking set of familiar compositions, even as the models themselves report ever-higher performance within that narrow band.

These three failure modes are not hypothetical; they are logical consequences of attention allocation operating without explicit safeguards. Once recognized, however, they become diagnosable. The conceptual framework developed here, therefore, equips the field to move beyond post-hoc lamentation of blind spots toward proactive design of attention-aware systems.

Table 2 formalizes the distinction between productive and pathological attention regimes, providing a comparative framework for diagnosing when model-driven focus accelerates discovery versus when it induces scientific myopia.

Table 2. Productive vs. pathological attention regimes in materials AI: a comparative framework

Dimension	Productive attention regime	Pathological attention regime	Underlying driver	Observable indicators	Long-term system effect
Exploration scope	Focused yet expandable subspace	Narrow, self-reinforcing subspace	Alignment vs. misalignment of objectives	Diversity of candidate materials explored	Expansion vs. contraction of chemical-space coverage
Model–data interaction	Balanced integration of new and diverse data	Reinforcement of existing data distributions	Data diversity vs. dataset bias	Rate of novel class inclusion	Knowledge diversification vs. stagnation
Active learning behavior	Strategic uncertainty reduction with boundary exploration	Exploitation-heavy acquisition cycles	Acquisition function design	Ratio of exploration vs. exploitation queries	Sustainable discovery vs. local optimization traps
Representation coverage	Inclusive of multiple material classes (ordered, disordered, hybrid)	Biased toward well-represented classes	Architectural inductive bias	Performance variance across material classes	Epistemic equity vs. representational injustice
Innovation trajectory	High-impact, non-linear breakthroughs	Incremental improvements in familiar systems	Attention diversification vs. concentration	Citation novelty and cross-domain synthesis	Transformative discovery vs. diminishing returns
Feedback dynamics	Controlled and periodically audited loops	Unchecked positive feedback loops	Governance vs. automation	Stability of attention	Adaptive system vs. lock-in and path dependence

				distribution over cycles	
--	--	--	--	-----------------------------	--

Scientific Attention Capital

Scientific attention capital is the finite, rivalrous, and path-dependent resource that ultimately limits the pace and direction of materials discovery. It comprises the total pool of researcher time, grant funding, synthesis and characterization capacity, peer-review bandwidth, and collective intellectual curiosity that the community can allocate at any given moment. Unlike computational FLOPs, which can be scaled indefinitely with additional hardware, scientific attention capital cannot be expanded without bound; every hour spent optimizing one class of materials is an hour unavailable for exploring another. In the attention economy of Materials AI, models do not generate this capital but act as its primary allocators, channeling it toward subspaces that the models have rendered legible, predictable, or high-yield while leaving other subspaces relatively starved. The introduction of this concept reframes Materials AI from a set of prediction tools into a sociotechnical governance layer that decides, often implicitly, how the community's scarcest resource will be invested across chemical space.

The capital metaphor is deliberate and generative. Just as financial capital compounds most rapidly when reinvested in already-liquid assets, scientific attention capital tends to accumulate in data-rich, model-familiar regions such as binary and ternary oxides or well-studied perovskites. Graph neural networks trained predominantly on crystalline structures [4, 21] learn attention weights that privilege local bonding motifs, thereby directing subsequent experimental campaigns and theoretical refinements toward those same motifs [1, 2, 25]. Active-learning pipelines further accelerate the compounding effect by selecting acquisition points that reduce uncertainty fastest within the model's current competence zone [9, 10, 18, 19]. Over multiple retraining cycles, the feedback loop converts an initial architectural bias into a durable stock of scientific attention capital: more publications appear, more grants are awarded, more students are trained, and the region becomes even more attractive for future investment. Meanwhile, sparsely populated regions—disordered alloys, flexible hybrid frameworks, or high-entropy ceramics—accumulate negative attention capital; their higher predictive uncertainty discourages allocation, which in turn keeps data density low and perpetuates their marginal status.

Who decides how this capital is spent? In the current paradigm, the decision is distributed across model architects, dataset curators, funding panels, and journal editors, yet the models themselves increasingly function as the de facto budget officers. When a transformer-based property predictor [13, 14] or a generative inverse-design engine [6, 7, 22] surfaces only the top-ranked candidates, it effectively performs a capital rationing step: the community's finite downstream resources are funneled into the narrow channel the model has highlighted. This delegation is efficient when the model's priorities align with genuine scientific needs, but it becomes problematic when the priorities reflect historical data artifacts rather than deliberate strategy. The attention economy, therefore, raises a governance question that the field has scarcely begun to address: should attention allocation remain an emergent property of performance-optimized training, or should it become an explicit, auditable design objective?

Objections and Replies

A natural objection to the attention-economy framing is that all scientific practice has always involved selective attention; models merely make an ancient feature of inquiry more visible and perhaps more powerful. While it is true that every experiment, every theoretical approximation, and every review article has historically pruned possibility space, the scale and automation introduced by Materials AI change the character of that selectivity. Human selectivity was episodic, distributed, and relatively transparent to the community; model-driven selectivity is continuous, centralized within surrogate functions, and largely opaque even to the model's own designers. When active-learning loops [9, 10, 18] operate across thousands of iterations without explicit human veto points, the cumulative redirection of scientific attention capital occurs at a speed and granularity that traditional peer-review mechanisms cannot readily correct. The objection,

therefore, underestimates the qualitative shift: what was once a manageable feature of individual cognition has become an infrastructural force shaping the entire discovery ecosystem.

A second objection asserts that models lack agency and thus cannot “shape” scientific attention in any meaningful sense; they are merely passive reflectors of the data and objectives supplied by human researchers. This view, while technically accurate at the level of individual forward passes, misses the emergent agential effects that arise in sociotechnical systems. The designers who choose graph-network inductive biases [4, 13, 21] or who define acquisition functions weighted toward uncertainty in specific subspaces [19, 20] embed value-laden choices that propagate through repeated deployment cycles. Once deployed at community scale, these choices acquire a quasi-independent momentum: the literature begins to cite the model's top candidates, funding calls reference the same subspaces, and new trainees orient their questions around the model's outputs. The models themselves do not possess intentions, yet they function as attention governors whose policies become self-reinforcing. Denying this emergent agency is analogous to claiming that recommendation algorithms on social platforms exert no influence simply because they are “just math”; the mathematics still reallocates collective attention with measurable consequences for what becomes salient.

A third common reply frames the entire discussion as merely a restatement of well-known dataset bias or model bias problems already addressed in the literature on high-throughput screening [11, 12] and feature importance [25]. Bias, however, is a static property of a dataset or a trained model; the attention-economy perspective reframes bias as a dynamic allocation process operating across time and across the research community. Bias explains why a model performs poorly on certain classes; attention capital explains why those classes then receive systematically less investment, fewer publications, and slower knowledge accumulation, further entrenching the bias. The distinction matters because remedies differ. Correcting static bias might involve reweighting loss functions or augmenting datasets; correcting pathological attention requires redesigning the feedback loops that govern how models ration scientific resources over successive generations of discovery. By moving beyond “bias” as a technical diagnostic to “attention capital” as a sociotechnical resource, the field gains a richer vocabulary for diagnosing and intervening in the long-term epistemic consequences of model deployment. These replies do not dismiss the objections but relocate them within a larger conceptual framework that treats attention allocation as an explicit, designable variable rather than an inevitable byproduct of doing science with machines.

Implications for Materials AI Practice

If the position developed here is taken seriously, Materials AI practice must shift from optimizing predictive accuracy in isolation to jointly optimizing accuracy and attention-allocation policy. Four actionable principles emerge as immediate consequences of viewing models as governors of scientific attention capital.

- Principle 1: Attention transparency by design. Every published Materials AI model should be accompanied by an explicit attention-allocation profile that reports not only internal architectural weights but also the downstream distribution of scientific attention capital that the model is likely to induce. For graph neural networks [4, 13, 21], this profile would quantify the fraction of chemical space that receives > 90% of the model's effective focus versus the fraction relegated to high-uncertainty fringes. Such profiles, analogous to carbon-footprint statements in computational chemistry, would allow reviewers and funding bodies to assess whether a new architecture risks pathological concentration or promotes productive diversification.
- Principle 2: Diversification regularization in training objectives. Current loss functions emphasize point-wise accuracy or uncertainty reduction; attention-aware training should incorporate explicit terms that penalize excessive concentration of attention capital within narrow compositional or structural families. Techniques already present in multi-task learning or uncertainty-aware active learning [9, 10, 18] can be extended with entropy-based regularizers that encourage models to maintain minimum attention investment across underrepresented regions, thereby counteracting the natural Matthew effect observed in data-rich chemistries [1, 2, 25-29].

- Principle 3: Periodic attention audits as a community norm. Just as high-throughput screening pipelines now require bias audits [11, 12], Materials AI deployments should undergo regular, independent attention audits that simulate the long-term redirection of scientific resources under realistic retraining scenarios. These audits would employ counterfactual chemical-space probes—regions deliberately held out from training—to quantify how the model's attention policy would allocate capital if deployed at scale. The resulting audit reports would be published alongside model weights, enabling the community to detect pathological patterns before they become entrenched.
- Principle 4: Hybrid human–model attention governance. Rather than treating active-learning loops as autonomous, attention-aware practice should embed mandatory human veto layers at key capital-allocation junctures—every 50th acquisition cycle, for example—where domain experts can explicitly redirect model focus toward scientifically motivated but model-undervalued subspaces. This hybrid governance preserves the efficiency gains of automation while restoring human agency over the ultimate distribution of scientific attention capital, ensuring that the attention economy remains subordinated to broader epistemic goals rather than dictating them.

Adopting these principles does not slow innovation; it channels it toward trajectories that are both more efficient and more plural. The field would move from an era in which attention allocation is an accidental side-effect of performance optimization to one in which it is a deliberate, auditable design criterion.

Conclusion

Materials AI systems are not neutral accelerators of discovery; they are attention-allocation infrastructures operating within a nascent attention economy of scientific research. By distinguishing architectural attention from scientific attention, by recognizing that models shape scientific attention capital through feature prioritization, search-space pruning, and output presentation, and by separating productive focus from pathological myopia, this perspective offers a conceptual framework for understanding how the very tools intended to expand possibility space can inadvertently contract it. The introduction of scientific attention capital as a finite, governable resource makes visible what was previously diffuse: the long-term epistemic consequences of architectural and training choices that have until now been evaluated almost exclusively on predictive metrics.

The position advanced here does not call for abandoning powerful architectures such as graph networks or active-learning strategies; it calls for explicit stewardship of the attention economy they govern. Implementing attention transparency, diversification regularization, periodic audits, and hybrid governance would transform Materials AI from an implicit shaper of scientific priorities into a deliberately designed mediator of them. The ultimate stakes are nothing less than the generative openness that has historically characterized materials science at its most creative moments. If the community embraces attention-aware design, the next decade of materials AI can expand the horizons of chemical possibility rather than narrowing them around the paths of least resistance. The call, therefore, is for systematic attention audits in every major materials AI system and for the integration of attention-allocation policy into the standard criteria of model evaluation, peer review, and funding allocation. Only by treating scientific attention capital as the scarce resource it truly is can the field ensure that its most powerful tools serve the broadest possible vision of discovery.

Acknowledgements

None

Conflict of interest

None

Financial Support

None

Ethics statement

None

Received: 26 Apr 2021

Revised: 10 Aug 2021

Accepted: 03 Sep 2021

Published online: 18 January 2022

Rights and permissions

Open Access The author(s) retain copyright. This article is licensed under the Creative Commons Attribution-NonCommercial-ShareAlike 4.0 International License. It may be shared and adapted for non-commercial purposes with appropriate attribution, an indication of changes, and distribution of adaptations under the same license. Third-party material may be subject to separate terms identified in its credit line. View the license at <https://creativecommons.org/licenses/by-nc-sa/4.0/>.

References

- Schmidt J, Marques MRG, Botti S, Marques MAL. Recent advances and applications of machine learning in solid-state materials science. *npj Comput Mater.* 2019;5:83.
<https://doi.org/10.1038/s41524-019-0221-0>.
- Butler KT, Davies DW, Cartwright H, Isayev O, Walsh A. Machine learning for molecular and materials science. *Nature.* 2018;559(7715):547-55.
<https://doi.org/10.1038/s41586-018-0337-2>.
- Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, et al. Attention is all you need. *Adv Neural Inf Process Syst.* 2017;30:5998-6008.
- 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.
<https://doi.org/10.1021/acs.chemmater.9b01294>.
- Nicol AA, Owens SM, Le Coze SS, MacIntyre A, Eastwood C. Comparison of high-technology active learning and low-technology active learning classrooms. *Act Learn High Educ.* 2018;19(3):253-65.
- Gómez-Bombarelli R, Wei JN, Duvenaud D, Hernández-Lobato JM, Sánchez-Lengeling B, Sheberla D, et al. Automatic chemical design using a data-driven continuous representation of molecules. *ACS Cent Sci.* 2018;4(2):268-76.
<https://doi.org/10.1021/acscentsci.7b00572>.
- Zunger A. Inverse design in search of materials with target functionalities. *Nat Rev Chem.* 2018;2(4).
<https://doi.org/10.1038/s41570-018-0121>.
- Wang AY, Kauwe SK, Murdock RJ, Sparks TD, Long S, Lutz RE, et al. Compositionally restricted attention-based network for materials property predictions. *npj Comput Mater.* 2021;7:77.
<https://doi.org/10.1038/s41524-021-00545-1>.

Lookman T, Balachandran PV, Xue D, Yuan R, Kirchdoerfer T, Gould T, et al. Active learning in materials science with emphasis on adaptive sampling using uncertainties for targeted design. *npj Comput Mater.* 2019;5:21.
<https://doi.org/10.1038/s41524-019-0153-8>.

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;11:5966.
<https://doi.org/10.1038/s41467-020-19597-w>.

Mazouze B, Nadon R, Makarenkov V. Identification and correction of spatial bias are essential for obtaining quality data in high-throughput screening technologies. *Sci Rep.* 2017;7:11921.
<https://doi.org/10.1038/s41598-017-11940-4>.

Caraus I, Mazouze B, Nadon R, Makarenkov V. Detecting and removing multiplicative spatial bias in high-throughput screening technologies. *Bioinformatics.* 2017;33(20):3258-67.

Louis SY, Zhao Y, Nasiri A, Wang X, Song Y, Liu F, et al. Graph convolutional neural networks with global attention for improved materials property prediction. *Phys Chem Chem Phys.* 2020;22(33):18141-8.
<https://doi.org/10.1039/D0CP01474E>.

Lin X, Jiang H, Wang L, Ren Y, Ma W, Zhan S. 3D-structure-attention graph neural network for crystals and materials. *Mol Phys.* 2022;120(11):e2051000.
<https://doi.org/10.1080/00268976.2022.2077258>.

Chen Q, Xiong Y, Chen X. Directed graph attention neural network utilizing 3D coordinates for molecular property prediction. *Comput Mater Sci.* 2021;200:110761.
<https://doi.org/10.1016/j.commatsci.2021.110761>.

Reiser P, Neubert M, Eberhard A, Torresi L, Zhou C, Shao C, et al. Graph neural networks for materials science and chemistry. *Commun Mater.* 2022;3:93.
<https://doi.org/10.1038/s43246-022-00315-6>.

Coley CW. Defining and exploring chemical spaces. *Trends Chem.* 2021;3(2):133-45.
<https://doi.org/10.1016/j.trechm.2020.12.001>.

Khalak Y, Tresadern G, Hahn DF, de Groot BL, Gapsys V. Chemical space exploration with active learning and alchemical free energies. *J Chem Theory Comput.* 2022;18(10):6259-70.
<https://doi.org/10.1021/acs.jctc.2c00752>.

Balachandran PV, Kowalski B, Sehirlioglu A, Lookman T, Cao Y, Ponomareva I, et al. Experimental search for high-temperature ferroelectric perovskites guided by two-step machine learning. *Nat Commun.* 2018;9:1668.
<https://doi.org/10.1038/s41467-018-03821-9>.

Yuan R, Liu Z, Balachandran PV, Xue D, Zhou Y, Ding X, et al. Accelerated discovery of large electrostrains in batio3-based piezoelectrics using active learning. *Adv Mater.* 2018;30(7).
<https://doi.org/10.1002/adma.201702884>.

Xie T, Grossman JC. Crystal graph convolutional neural networks for an accurate and interpretable prediction of material properties. *Phys Rev Lett.* 2018;120(14):145301.
<https://doi.org/10.1103/PhysRevLett.120.145301>.

Sanchez-Lengeling B, Aspuru-Guzik A. Inverse molecular design using machine learning: Generative models for matter engineering. *Science.* 2018;361(6400):360-5.

Schütt KT, Sauceda HE, Kindermans PJ, Tkatchenko A, Müller KR. SchNet: A deep learning architecture for molecules and materials. *J Chem Phys.* 2018;148(24):241722.

<https://doi.org/10.1063/1.5019779>.

Isayev O, Oses C, Toher C, Gossett E, Curtarolo S, Tropsha A, et al. Universal fragment descriptors for predicting properties of inorganic crystals. *Nat Commun.* 2017;8:15679.

<https://doi.org/10.1038/ncomms15679>.

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

<https://doi.org/10.1038/s41524-017-0056-5>.

Himanen L, Geurts A, Foster AS, Rinke P. Erratum: Data-driven materials science: Status, challenges, and perspectives. *Adv Sci (Weinh).* 2020;7(2):1903667. Erratum for: *Adv Sci (Weinh).* 2019;6(21):1900808.

Tshitoyan V, Dagdelen J, Weston L, Dunn A, Rong Z, Kononova O, et al. Unsupervised word embeddings capture latent knowledge from materials science literature. *Nature.* 2019;571(7763):95-8.

<https://doi.org/10.1038/s41586-019-1335-8>.

Jha D, Ward L, Paul A, Liao WK, Choudhary A, Wolverton C, et al. ElemNet: Deep learning the chemistry of materials from only elemental composition. *Sci Rep.* 2018;8(1):17593.

<https://doi.org/10.1038/s41598-018-35934-y>.

Pilania G, Gubernatis JE, Lookman T. Multi-fidelity machine learning models for accurate bandgap predictions of solids. *Comput Mater Sci.* 2017;129:156-63.

<https://doi.org/10.1016/j.commatsci.2016.12.028>.