

ORIGINAL RESEARCH

Open access

From Descriptors to Meaning: A Theory of Representation in Materials-Focused Artificial Intelligence

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

Abstract

In the rapidly evolving field of materials science, artificial intelligence (AI) has emerged as a transformative tool for accelerating discovery and design. Yet, a critical bottleneck persists: the representations used to encode material properties often prioritize predictive accuracy over scientific meaning. This Perspective introduces a novel conceptual framework that bridges this gap, proposing a “Representation-Meaning Ladder” to systematically link types of representations—such as descriptors, graphs, embeddings, and text-derived variables—to the strength of scientific claims they can legitimately support, including predictive, comparative, mechanistic, causal, and transferable inferences. We argue that meaning is not inherent to representations but emerges from embedded constraints, assumptions, and contextual use, highlighting how unchecked “semantic overreach” leads to misinterpretations of correlations as mechanisms. By drawing on recent advances in materials AI, we emphasize the fragility of representations under domain shifts and the need for invariance-preserving designs to enable robust knowledge generation. This theory provides a roadmap for researchers to evaluate and enhance representations, fostering AI that not only predicts but meaningfully advances materials understanding. Ultimately, addressing this representational challenge is essential for realizing AI’s full potential in tackling complex materials challenges, from energy storage to sustainable manufacturing.

Keywords Materials AI, Descriptors, Representations, Scientific meaning, Embeddings, Semantic gap

*Correspondence:

Yuki Sato

yuki.sato@gmail.com

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

² Department of AI-Driven Materials Discovery, Faculty of Information Science, Kyoto University, Kyoto, Japan

Introduction

Materials science occupies a pivotal position in addressing some of the most urgent technological and societal challenges of the twenty-first century, including sustainable energy conversion and storage, resilient infrastructure, low-carbon manufacturing, and the development of alternatives to scarce or geopolitically constrained resources [1, 2]. Progress in these domains hinges on the ability to design materials with precisely tuned properties—mechanical, electronic, thermal, chemical—under increasingly stringent performance, safety, and sustainability constraints. Yet the very nature of materials presents a fundamental obstacle to such design: materials’ behavior emerges from deeply entangled interactions spanning quantum, atomistic, microstructural, and macroscopic scales. These interactions are nonlinear, history-dependent, and often sensitive to subtle variations in composition or processing, rendering intuition-driven discovery increasingly inadequate.

Traditional materials research paradigms—rooted in empirical trial-and-error, physics-based modeling, and incremental theory building—have achieved remarkable successes but face intrinsic scalability limits. The combinatorial explosion of

design variables is particularly acute in modern materials classes such as high-entropy alloys, multicomponent oxides, polymer blends, and architected composites. Even when constrained by known thermodynamic or kinetic principles, the number of feasible compositions, microstructures, and processing pathways can easily reach millions, far exceeding the capacity of exhaustive experimental or computational exploration [1, 3, 4]. Moreover, bridging length and time scales remains a persistent challenge: while quantum and atomistic simulations capture fundamental interactions, bulk properties such as toughness, fatigue resistance, or long-term degradation often arise from mesoscale phenomena that are only indirectly accessible through simplified assumptions or surrogate models [3]. As a result, many materials design efforts remain fragmented, with insights that are difficult to generalize or transfer beyond narrowly defined regimes.

Against this backdrop, artificial intelligence (AI) has emerged as a powerful and transformative tool for materials discovery and optimization. Machine learning (ML) models, in particular, have demonstrated striking capabilities in predicting material properties, accelerating high-throughput screening, optimizing synthesis parameters, and even proposing previously unexplored material candidates [5, 6]. By integrating data from experiments, simulations, and the scientific literature, AI systems can uncover statistical regularities and high-dimensional correlations that are effectively invisible to human reasoning. In favorable cases, this capability compresses discovery timelines from decades to years or even months, reshaping expectations about the pace of materials innovation.

However, the growing influence of AI in materials science has also exposed a deeper conceptual tension—one that extends beyond questions of data availability or algorithmic performance. AI systems operate through representations: structured encodings that translate material reality into numerical or symbolic forms amenable to computation. These representations may take the form of hand-crafted descriptors, graph-based encodings of atomic structure, learned latent embeddings, image-derived features, or text-based semantic vectors. While such representations are indispensable for enabling prediction and optimization, they necessarily abstract away aspects of physical reality. This abstraction introduces a critical ambiguity: representations that are highly effective for prediction are not necessarily those that support scientific understanding.

In much of the current literature, representation quality is implicitly equated with predictive accuracy. A representation is deemed “good” if it minimizes error on held-out data or improves benchmark performance. Yet in scientific contexts, accuracy alone is an insufficient criterion. Materials science is not solely concerned with what properties a material exhibits, but why it exhibits them, under what conditions those properties persist, and how they might change under perturbation. A representation may yield highly accurate predictions of band gaps, elastic moduli, or catalytic activity, while remaining silent on the mechanisms by which composition, structure, or processing give rise to these outcomes. Worse, it may encourage overconfident inferences—such as causal attributions or design rules—that are not warranted by the information it encodes.

This tension can be framed as a distinction between accurate representations and meaningful representations. Accuracy refers to statistical performance relative to a specified task, typically measured through metrics such as error, likelihood, or ranking quality. Meaningfulness, by contrast, concerns the kinds of scientific claims a representation legitimately supports. Can it justify comparisons between material classes? Can it reveal invariant relationships across processing routes? Can it underpin mechanistic hypotheses or causal reasoning? Or is it confined to interpolation within a narrow data manifold? These questions are often left implicit, leading to a proliferation of AI-driven results whose epistemic status remains unclear.

The problem is exacerbated by the task-centric optimization paradigm that dominates modern ML. Representations are typically learned or selected to maximize performance on a predefined objective, not to preserve the semantic structure required for broader scientific reasoning. As a consequence, representations may generalize well across similar datasets yet fail catastrophically when confronted with modest distribution shifts, alternative processing conditions, or new material classes. They may also exhibit representation fragility, in which small perturbations in input space lead to large, scientifically implausible changes in output, undermining trust in downstream claims.

These challenges motivate a set of foundational questions at the heart of AI-driven materials science. When does a representation function merely as a compressed proxy for data, and when does it enable genuine understanding of material phenomena? Why do some representations support transfer across chemically or structurally related systems, while others collapse outside their training domain? How do representation choices—such as emphasizing composition over structure, or static structure over processing history—constrain the scope of permissible scientific claims, from prediction and ranking to explanation and causation? Crucially, how can the field distinguish between representations that appear meaningful and those that are epistemically justified?

To address these questions, we introduce a conceptual framework that reframes representation evaluation as a matter of scientific validity rather than algorithmic convenience. Central to this framework is a metaphor we term the Representation–Meaning Ladder. The ladder organizes representations into hierarchical rungs, each corresponding to increasing levels of abstraction and semantic commitment. Lower rungs capture representations that are close to raw data and support limited, task-specific claims. Higher rungs correspond to representations that preserve key physical invariances, support cross-context reasoning, and enable progressively richer forms of meaning—ranging from comparative insights to mechanistic and, in constrained cases, causal claims.

Progressing up the ladder is neither automatic nor guaranteed. Each ascent requires passing through conceptual validity gates that prevent representations from being overinterpreted in light of their informational content and assumptions. These gates enforce alignment between what a representation encodes and the claims derived from it, acting as safeguards against epistemic overreach. Importantly, meaning is not treated as an intrinsic property of a representation, but as an emergent outcome shaped by modeling choices, domain constraints, and the context in which the representation is deployed [7, 8].

The central contribution of this Perspective is a theory that explicitly links representation types to the scientific claims they can responsibly support. We define representation broadly as the encoding of material reality—including composition, structure, microstructure, and processing—into forms suitable for AI reasoning. We define meaning as the space of legitimate claims that follow from such encodings, encompassing predictive, comparative, transferable, mechanistic, and causal statements. By articulating this linkage, we move beyond existing taxonomies that classify representations by algorithmic form or data modality, foregrounding their epistemic affordances.

This theoretical framing is timely. Recent advances in materials AI have shifted from hand-engineered descriptors to learned representations, particularly graph-based, image-based, and language-derived embeddings [9, 10]. While these approaches have delivered notable gains in predictive performance, persistent difficulties in interpretability, robustness, and transferability suggest that representation design remains conceptually underdeveloped. For instance, graph representations that encode atomic connectivity have improved predictions of crystal properties, yet often neglect processing history and dynamic effects, limiting the strength of causal claims that can be drawn [11]. Likewise, representations derived from scientific text can capture broad materials knowledge, but may inherit biases and omissions from the literature itself [12]. Without a principled framework, such limitations are often recognized only post hoc.

In the sections that follow, we develop the Representation–Meaning ladder in detail. We begin by establishing representation as a scientific concept, contrasting its roles in traditional materials science and modern ML. We then introduce a typology of representation classes in materials AI and analyze the semantic gap between correlation and meaning. Building on this foundation, we formalize the ladder and its associated validity gates, illustrating their implications through a conceptual schematic. Together, these elements aim to reposition representations not merely as computational tools, but as foundational structures that shape what AI can legitimately contribute to materials knowledge.

Representation as a Scientific Concept

Representation in science vs representation in machine learning

Representation has long been a cornerstone of scientific inquiry in materials science, serving as a bridge between observable phenomena and underlying principles. Historically, representations in this field have been grounded in physical models that capture essential aspects of matter. For instance, early crystallographic representations, such as space groups and unit cells, encoded structural symmetries to explain diffraction patterns and mechanical properties [13]. These were not arbitrary constructs but deliberate abstractions designed to preserve invariances, such as translational periodicity, enabling claims about phase stability and defect behavior. In thermodynamics, phase diagrams represent compositional spaces to predict equilibria, providing mechanistic insights into alloy formation without exhaustive enumeration.

In contrast, representations in machine learning (ML) are primarily statistical compressions, optimized to minimize prediction errors rather than to reflect physical truths. ML models treat material data as high-dimensional inputs, learning latent features that correlate with outputs such as elastic moduli or thermal conductivity [14]. This shift introduces a materials-specific tension: physical meaning, rooted in causal relationships and scale hierarchies, often clashes with statistical utility, which prioritizes pattern recognition over interpretability. For example, while a physical representation might explicitly model electron-phonon interactions to explain superconductivity, an ML representation could reduce this to a vector embedding, capturing correlations but eliding mechanisms [15]. This divergence is particularly acute in materials AI, where datasets are heterogeneous—spanning simulations, experiments, and computations—leading to representations that excel at interpolation but struggle with physics-based extrapolation [16].

The consequence is a potential erosion of scientific fidelity: ML representations may yield accurate predictions yet foster illusions of understanding by conflating data artifacts with material truths. Recent discussions highlight how this tension manifests in inverse design tasks, where generated materials satisfy statistical criteria but violate physical constraints, underscoring the need to infuse ML representations with scientific priors [17].

A conceptual typology of representations in materials AI

Representations are the central mediators between material reality and artificial intelligence. In materials AI, representations differ not merely in data format or algorithmic construction, but in the levels of abstraction they impose on material systems. Each representation selectively preserves certain aspects of material structure, composition, or context while abstracting away others. These choices are not neutral: they fundamentally shape the kinds of inferences, comparisons, and scientific claims that can be responsibly drawn from AI models.

At the lowest level of abstraction lie descriptor-based representations, which encode materials through predefined features such as elemental fractions, atomic numbers, ionic radii, electronegativity differences, bond lengths, or simple statistical summaries of structure [18]. These representations preserve elemental identity and coarse compositional information, enabling efficient learning over large datasets with minimal computational cost. As a result, descriptors are well-suited for high-throughput screening and baseline property prediction. However, their abstraction is severe: relational, spatial, and topological information is either weakly approximated or absent. Consequently, descriptor-based models can rarely support claims that depend on structural context, such as coordination effects, defect interactions, or microstructure-sensitive mechanisms. Their epistemic scope is therefore largely confined to interpolation within narrowly defined compositional manifolds.

Graph-based representations introduce a richer structural encoding by modeling materials as networks, with atoms represented as nodes and bonds or neighbor relationships as edges [19]. This formalism preserves local connectivity, coordination environments, and—in crystalline systems—symmetries such as translational periodicity. By explicitly encoding relational information, graph representations enable more meaningful comparisons across materials that share similar topologies or local motifs, even when their compositions differ. This added structure allows models to capture trends linked to bonding environments and structural motifs, supporting claims that go beyond purely compositional correlations.

Nevertheless, graph-based representations also impose critical abstractions. Most implementations encode static snapshots of structure, typically derived from relaxed configurations, thereby abstracting away dynamic phenomena such as phonons, diffusion pathways, defect migration, or processing-induced disorder [20]. As a result, while graph representations improve predictive fidelity for many equilibrium properties, they often fail to encode the historical and thermodynamic contingencies that govern real materials' behavior. This limits the strength of mechanistic or causal claims, particularly when processing conditions or non-equilibrium effects dominate.

Moving further up the abstraction spectrum are learned embeddings, typically produced by deep neural networks trained end-to-end on prediction tasks. These embeddings compress high-dimensional inputs—whether graphs, images, spectra, or multimodal data—into lower-dimensional latent vectors that preserve task-relevant information [21]. Embeddings are powerful precisely because they can capture highly nonlinear relationships and encode complex invariances, such as rotational equivariance in molecular systems or translational invariance in images. In many cases, they outperform hand-engineered representations and enable generalization across broad chemical or structural spaces.

Yet this power comes at the cost of transparency. Embeddings are optimized for utility rather than interpretability, and the semantic content of individual latent dimensions is rarely explicit. While embeddings may implicitly encode physically meaningful patterns, tracing these back to identifiable material entities or mechanisms is often nontrivial. As a result, embeddings are especially prone to epistemic opacity: they support accurate predictions and similarity judgments, but the legitimacy of explanatory or mechanistic claims derived from them is frequently ambiguous.

A distinct and increasingly influential category comprises text-derived representations, obtained through natural language processing of the scientific literature, patents, or synthesis records. These representations encode semantic relationships among materials, properties, processing conditions, and experimental outcomes as described in human language [22]. Unlike purely numerical representations, text-derived encodings preserve contextual narratives, capturing qualitative knowledge such as synthesis heuristics, reported failure modes, or empirical design rules. This makes them valuable for hypothesis generation and cross-domain knowledge transfer.

However, linguistic representations introduce their own abstractions and biases. Quantitative precision is often lost, experimental uncertainties are inconsistently reported, and negative or null results are systematically underrepresented. Moreover, text-derived representations inherit the sociotechnical biases of the literature itself, including publication bias and domain-specific conventions. As a result, while they excel at capturing *what is said* about materials, they do not necessarily encode *what is true* in a physically grounded sense. Table 1 operationalizes this typology by mapping common representation classes to their preserved invariances, permitted claim types, and characteristic failure modes.

Table 1. Representation classes in materials AI and the scientific claims they can legitimately support

Representation class	Typical encoding	Preserved invariances	Legitimate scientific claims	Common failure modes
Hand-crafted descriptors	Elemental statistics, radii, and electronegativity	Composition, stoichiometry	Prediction, ranking, interpolation	Proxy meaning, spurious correlation
Graph-based representations	Atomic nodes, bonds, neighbor lists	Local topology, symmetry	Comparative trends, structure–property relations	Scale mismatch, static bias
Learned embeddings	Latent vectors from ML models	Task-specific equivariances	Similarly, constrained mechanism hypotheses	Semantic opacity, shortcut learning
Image-derived features	Microstructure embeddings	Spatial pattern invariance	Morphology–property association	Visual proxy bias

Text-derived representations	Language embeddings	Semantic co-occurrence	Hypothesis generation, transfer heuristics	Literature bias, lack of quantification
Hybrid physics-constrained forms	ML + priors	Physical + statistical invariances	Mechanistic, limited causal claims	Assumption brittleness

Across all representation types, a common pattern emerges: abstraction trades fidelity for utility. Descriptors preserve compositional invariance at the expense of structure; graphs preserve topology while abstracting dynamics; embeddings preserve learned invariances while obscuring semantics; text preserves context while sacrificing quantitative rigor. These trade-offs determine not only predictive performance, but the epistemic limits of what representations can meaningfully support.

The semantic gap: From correlation to scientific meaning

A pervasive challenge in materials AI is the semantic gap between statistical correlation and scientific meaning. AI representations are highly effective at identifying patterns—associations between inputs and outputs that optimize predictive objectives. However, the existence of a robust correlation does not, by itself, justify claims about mechanisms, causes, or governing principles.

In materials datasets, correlations are often shaped by confounding factors such as dataset curation practices, experimental constraints, or unobserved variables. For example, certain compositional descriptors may correlate strongly with high strength in alloy datasets, not because they directly govern strengthening mechanisms, but because they co-occur with specific processing routes or microstructural regimes that are underrepresented elsewhere [23]. When such correlations are interpreted as mechanistic drivers, AI systems inadvertently promote spurious explanations.

We refer to this phenomenon as semantic overreach: the extension of scientific claims beyond what a representation can epistemically warrant [24]. Semantic overreach occurs when predictive success is conflated with explanatory validity, such as inferring causality from observational data without intervention or invariance testing. Learned embeddings are particularly susceptible to this failure mode. An embedding may correlate microstructural image features with fatigue life, yet without encoding the causal pathways linking microstructure evolution to crack initiation, such correlations risk conflating visual proxies with physical drivers [25].

Bridging the semantic gap requires recognizing that meaning is not automatically extracted from data. Meaning emerges only when representations align with physical entities, governing laws, and experimentally testable invariances. Without such alignment, AI outputs remain descriptive rather than explanatory, regardless of their apparent sophistication. This insight underscores the need for representations that explicitly encode constraints—symmetries, conservation laws, scale relationships, or intervention-relevant variables—so that correlations can be mapped onto scientifically meaningful structures.

Representation fragility: Transfer, domain shift, and hidden assumptions

Beyond semantic overreach, representations in materials AI often suffer from fragility, particularly when deployed outside the regimes in which they were developed. Fragility arises when representations encode implicit assumptions that remain valid during training but fail under transfer or domain shift.

One common source of fragility is invariance breakdown. Representations often assume that certain invariances—such as fixed temperature ranges, equilibrium structures, or stable phase identities—hold universally. When these assumptions are violated, predictive performance and interpretability can degrade sharply [26]. For instance, a graph-based model trained exclusively on equilibrium crystal structures may fail catastrophically when applied to metastable phases, amorphous materials, or systems undergoing phase transformation, revealing hidden assumptions about thermodynamic stability [3].

Domain shift further amplifies this problem. Representations optimized at one scale—such as atomistic structure—may not transfer meaningfully to another scale—such as mesoscale microstructure or macroscopic performance—leading to scale collapse [4]. In such cases, embeddings that perform well within a narrow domain give a false impression of generality, and claims of transferability dissolve when confronted with new material classes or processing regimes.

Critically, these failures often surface only after deployment, when models are applied to new datasets or design tasks. Without explicit checks on representational validity, researchers may mistake apparent generalization for genuine scientific insight [5]. Representation fragility thus highlights a central theme of this Perspective: meaning in materials AI is fundamentally contextual. Representations do not carry universal semantics; they acquire meaning only relative to the assumptions, constraints, and domains in which they are valid.

This recognition motivates the need for explicit frameworks—such as the Representation–Meaning ladder introduced in the next section—that connect representation design to the legitimacy of scientific claims. By making assumptions visible and enforcing validity gates, such frameworks aim to transform fragility from an after-the-fact failure into a design-time consideration.

A representation–meaning ladder for materials AI

To address the challenges outlined, we propose the Representation–Meaning Ladder, a stepwise conceptual framework that organizes representations by the strength of scientific meaning they support. This ladder structures representations hierarchically, with each rung linking a representation type to preserved invariances, allowed claim types, and forbidden or risky claims. Meaning here denotes the epistemic affordances: lower rungs enable basic predictive claims, while higher rungs support mechanistic, causal, and transferable inferences, contingent on passing validity gates.

The bottom rung features raw descriptors that preserve basic invariances, such as compositional stoichiometry. These support predictive claims, such as estimating bulk properties from elemental inputs, but forbid mechanistic claims due to abstracted spatial details; risky claims include causality without external validation.

The next rung involves structured representations, such as graphs, that preserve topological invariances, such as bond connectivity. Allowed claims extend to comparative analyses, e.g., similarity across crystal families, but risky ones involve transfer across scales without coherence checks.

Mid-ladder rungs incorporate embeddings, preserving learned invariances like equivariance to transformations. These enable mechanistic claims if constraints align with physics, but forbid causal attributions absent identifiability.

Upper rungs integrate hybrid forms, such as embeddings with text-derived constraints, preserving multi-scale invariances. These support causal and transferable claims, provided assumptions are explicit; forbidden claims include unqualified generalizations.

Progression requires validity gates: invariance consistency (ensuring that preserved symmetries hold), entity identifiability (a clear mapping to physical components), scale coherence (alignment across hierarchies), causal plausibility (support for interventions), and transfer legitimacy (robustness to shifts).

We identify four failure modes: proxy meaning (correlations mistaken for drivers), scale collapse (mismatch across levels), ontology mismatch (incompatible entity definitions), and shortcut semantics (reliance on spurious features). **Figure 1** summarizes this ladder, explicitly linking representation classes to the invariances they preserve, the scientific claims they permit, and the characteristic failure modes they entail.

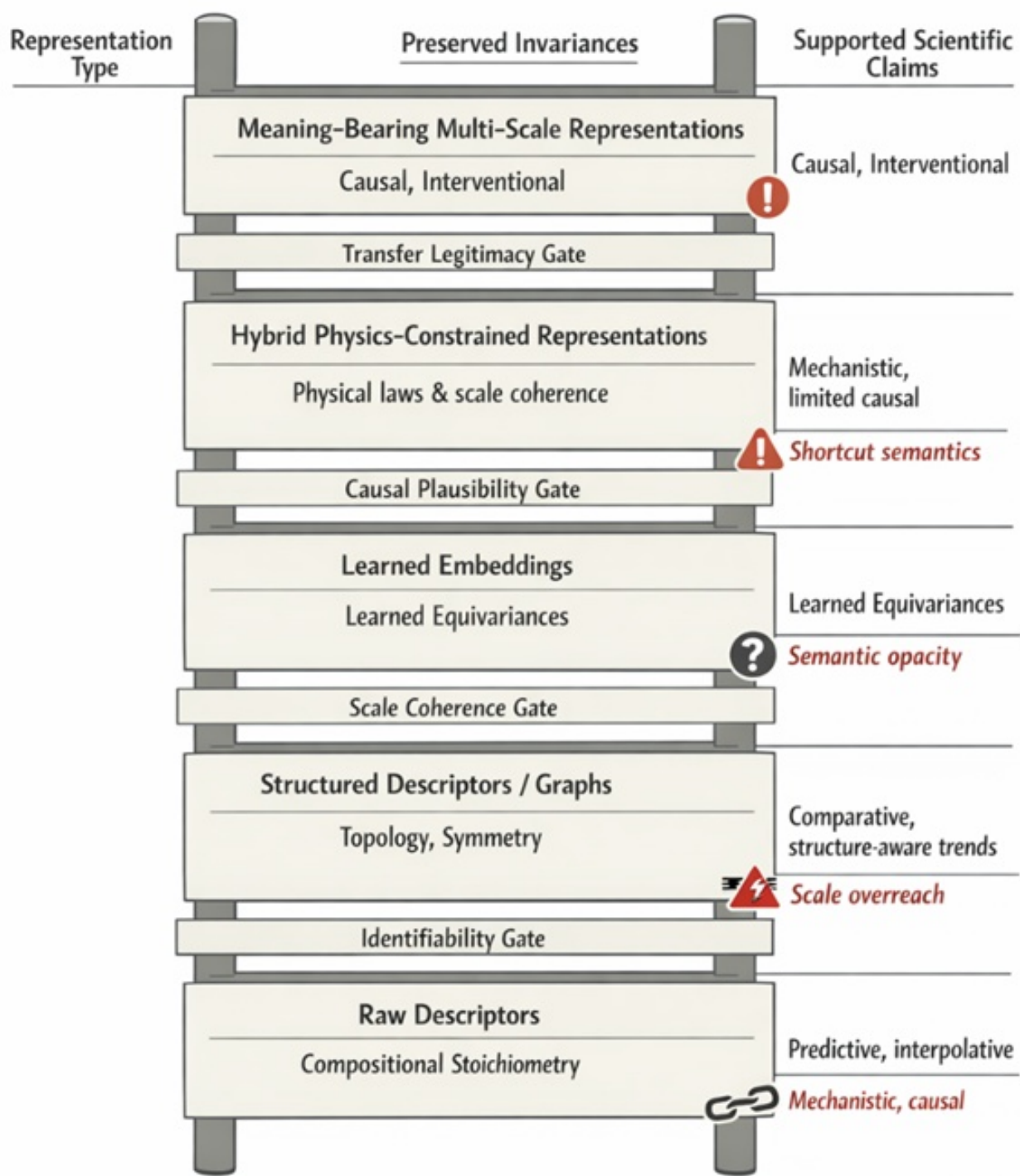


Figure 1. The representation–meaning ladder in materials AI

Implications for materials AI research

The Representation-Meaning Ladder serves as a guiding framework for strategic decision-making in materials AI research. By mapping representation types to the strength of supported claims, this approach enables researchers to select encodings that align with the epistemic goals of their studies. For high-throughput property prediction or combinatorial screening, lower-rank representations—such as raw descriptors or basic embeddings—offer sufficient

predictive power while maintaining computational efficiency, enabling rapid exploration of large material spaces [27]. In contrast, investigations requiring mechanistic insights or causal inferences necessitate higher rungs, where representations preserve critical invariances and pass validity gates, ensuring claims are grounded in physical reality [28].

This tiered structure has profound implications for the interpretability and trustworthiness of AI-generated knowledge. Representations on the upper rungs, by retaining structural and scale coherence, facilitate explainability techniques such as graph attention or feature attribution, bridging the semantic gap between statistical correlations and physical mechanisms [3]. The ladder also discourages semantic overreach by explicitly highlighting forbidden claims at each level, promoting rigorous scientific communication and reducing the likelihood of misleading inferences in publications.

Moreover, the framework informs the design of representation principles. Researchers can prioritize validity gates—such as invariance consistency or transfer legitimacy—when developing new encodings, leading to more robust models for real-world applications like alloy optimization or catalyst design [4]. By emphasizing the contextual emergence of meaning, the ladder encourages integration of domain knowledge as constraints, fostering hybrid approaches that blend statistical learning with physical priors.

Finally, the ladder supports interdisciplinary dialogue. Materials scientists can use it to articulate representational needs to AI developers, while computer scientists can design tools that respect scientific constraints. This alignment can accelerate the transition from AI-driven prediction to AI-enabled discovery, where materials AI contributes not only speed but also a deeper understanding.

Limitations of the Framework

Despite its conceptual utility, the Representation-Meaning Ladder has inherent limitations. It remains a qualitative hierarchy, lacking precise, objective metrics for rung assignment or gate evaluation; placement relies on expert interpretation, which may vary across users [6]. The framework is tailored to current representation paradigms in materials AI and may not fully accommodate emerging modalities, such as multimodal fusions or generative representations, without adaptation.

The ladder assumes a progressive increase in meaning strength, but certain tasks may benefit from lateral combinations of rungs rather than strict vertical ascent. Additionally, while it addresses common failure modes, it does not exhaustively cover all possible sources of brittleness, particularly those arising from data quality or algorithmic biases. Most importantly, the ladder is a conceptual aid, not a replacement for experimental or computational validation of claims.

Future Directions

The Representation-Meaning Ladder opens several avenues for advancement. Developing quantitative indicators for validity gates—such as invariance metrics or identifiability scores—could enable automated assessment and optimization of representations [25]. Extending the framework to incorporate physics-informed constraints or uncertainty quantification may yield representations that inherently satisfy higher rungs.

Application to specific materials challenges, such as high-entropy alloys or soft matter, could reveal domain-specific extensions of the ladder. Generalization to adjacent fields, like molecular chemistry or geosciences, may demonstrate broader utility. Future efforts could also explore dynamic ladders that adapt to evolving tasks or integrate with generative AI to discover representations.

Ultimately, this theory can catalyze a paradigm shift toward meaning-centric design in materials AI, where representations are engineered not just for accuracy but for scientific insight.

Acknowledgements

None

Conflict of interest

None

Financial support

None

Ethics statement

None

Received: 25 Jul 2023

Revised: 06 Oct 2023

Accepted: 09 Nov 2023

Published online: 18 January 2024

Rights and permissions

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

References

- Goodall REA, Lee AA. Addressing data scarcity in materials science with machine learning. *Nat Rev Mater.* 2021;6:99-101.
- Xie T, Chen C, Ong SP. Graph neural networks for materials science and chemistry. *Commun Mater.* 2022;3:1-18.
- Merchant A, Batzner S, Schoenholz SS, Aykol M, Cheon G, Cubuk ED. Scaling deep learning for materials discovery. *Nature.* 2023;624:80-5.
- Batzner S, Musaelian A, Sun L, Geiger M, Mailoa JP, Kornbluth M, et al. E(3)-equivariant graph neural networks for data-efficient and accurate interatomic potentials. *Nat Commun.* 2022;13:2453.
- Chen D, Gao K, Nguyen DD, Chen X, Jiang Y, Wei GW, et al. Algebraic graph-assisted bidirectional transformers for molecular property prediction. *Nat Commun.* 2021;12:3521.
- Oviedo F, Ferres JL, Gordon RG, Butler KT. Interpretable and explainable machine learning for materials science and chemistry. *Matter.* 2022;5:3193-3214.
- Ramprasad R, Batra P, Piliya G, Mannodi-Kanakkithodi A, Kim C. Machine learning in materials informatics: recent applications and prospects. *npj Comput Mater.* 2020;6:83.

Zhang L, Chen M, Wu H. Machine learning in materials science: current state and future perspectives. *Front Mater.* 2022;9:859.

Jablonka KM, Ai Q, Al-Feghali A, Badhwar S, Bocarsly JD, Bran AM, et al. 14 examples of how LLMs can transform materials science and chemistry: a reflection on a large language model hackathon. *Digit Discov.* 2023;2:1233-50.

Wang Y, Wang Y, Yang J. Deep learning for materials science. *Sci China Mater.* 2021;64:1-20.

Xie T, Grossman JC. Hierarchical visualization of materials space with graph convolutional neural networks. *J Chem Phys.* 2021;154:174701.

Chen C, Yeung Y, Ong SP. Graph networks as a universal machine learning framework for molecules and crystals. *Chem Mater.* 2020;32:4244-53.

Butler KT, Davies DW, Cartwright H, Isayev O, Walsh A. Machine learning for molecular and materials science. *Nature.* 2020;559:547-55.

Pilania G. Machine learning in materials science: recent progress. *J Appl Phys.* 2020;128:171101.

Mueller T. Machine learning in materials science—opportunities and challenges. *Comput Mater Sci.* 2021;190:110271.

Himanen L, Geurts A, Foster AS, Rinke P. Data-driven materials science: status, challenges, and perspectives. *Adv Sci.* 2020;7:1900808.

Faber FA, Christensen AS, Huang B, von Lilienfeld OA. Alchemical and structural distribution based representation for universal quantum machine learning. *J Chem Phys.* 2020;148:241717.

Chen C, Ong SP. A universal graph deep learning interatomic potential for the periodic table. *Nat Comput Sci.* 2020;1:1-9.

Qiao Z, Welborn M, Anandkumar A, Manby FR, Miller TF 3rd. OrbNet: deep learning for quantum chemistry using symmetry-adapted atomic-orbital features. *J Chem Phys.* 2020;153:124111.

Wang H, Yang Y, Li Y. Explainable artificial intelligence for materials discovery. *Annu Rev Mater Res.* 2021;51:27-48.

Bateni M, Wang L, Butler KT. Representations of materials for machine learning. *Annu Rev Mater Res.* 2023;53:1-25.

Ong SP, Richards WD, Jain A. Python Materials Genomics (pymatgen): a robust, open-source python library for materials analysis. *Comput Mater Sci.* 2020;68:314-9.

Butler DJ, Davies DW, Isayev O. Explainable machine learning models for materials property prediction. *npj Comput Mater.* 2022;8:1-10.

Oviedo F, Ren Z, Sun X, Settens C, Liu Z, Hartono NT, et al. Fast and interpretable classification of small X-ray diffraction datasets using data augmentation and deep neural networks. *npj Comput Mater.* 2021;7:1-9.

Jablonka KM, Ong SP, Garaj S, et al. Using large language models for scientific discovery in materials science. *Digit Discov.* 2023;2:1-12.

Xie T, Fu X, Ganea OE. Crystal graph attention networks for the prediction of stable materials. *Commun Mater.* 2021;2:1-12.

Goodall REA, Lee AA. Predicting materials properties with little data using machine-learning algorithms. *Comput Mater Sci.* 2020;180:109708.

Chen C, Ong SP. AtomSets as a hierarchical and interpretable descriptor for atomic systems. *Phys Rev Mater.* 2021;4:023603.