

ORIGINAL RESEARCH

Open access

Uncertainty as a Design Signal: A Conceptual View of Confidence, Risk, and Action in Materials AI

Ravi Kumar^{1*}, Neha Sharma¹, Aniket Deshmukh², Arjun Nair¹

Abstract

In the domain of applied artificial intelligence (AI) for materials science, uncertainty emerges as a pivotal signal that informs design decisions, yet its conceptual interpretation remains underexplored. This paper delineates uncertainty as the lack of complete knowledge about a system's state or outcomes, distinct from confidence, which reflects a model's self-assessed reliability in predictions; risk, which weights uncertainty by potential consequences; and actionability, which denotes the warrant for proceeding with design actions based on interpreted signals. Traditional approaches often conflate these concepts, leading to suboptimal decisions in materials discovery and optimization. For instance, high confidence in AI predictions may not equate to low risk in high-stakes applications like alloy design for extreme environments, where epistemic gaps could amplify failures. This conceptual manuscript proposes a decision-theoretic framework, the Uncertainty-to-Action Map, that translates uncertainty types—epistemic, aleatory, and semantic—into risk postures and subsequent action classes, such as screening candidates, prioritizing explorations, deferring judgments, redesigning models, stopping pursuits, hedging bets, or diversifying portfolios. By incorporating gates for stake assessment, ambiguity detection, domain scope evaluation, cost asymmetry analysis, and stopping logic, the framework mitigates failure modes like overconfidence and decision paralysis. This model fosters a nuanced view, emphasizing that uncertainty, when properly interpreted, serves as a design asset rather than a hindrance, promoting robust AI-assisted materials innovation.

Keywords Materials AI, Uncertainty interpretation, Actionability, Decision theory, Risk posture, Confidence calibration

*Correspondence:

Ravi Kumar

ravi.kumar@gmail.com

¹ Department of Materials Engineering and Data Science, Faculty of Engineering, IIT Delhi, New Delhi, India

² Department of AI-Based Materials Design, Faculty of Engineering, IIT Bombay, Mumbai, India

Introduction

The integration of artificial intelligence (AI) into materials science has transformed the landscape of design and discovery, enabling accelerated exploration of vast chemical spaces and more accurate predictions of complex properties. However, this advancement introduces profound challenges related to uncertainty, which must be interpreted as a design signal rather than dismissed as mere noise [1–4]. In materials AI, uncertainty arises from inherent limitations in data, models, and extrapolations, requiring a conceptual shift to view it as an informative cue for guiding actions. This interpretation is essential because raw uncertainty outputs from AI systems—such as variance in ensemble predictions or Bayesian posteriors—do not inherently dictate design pathways; they require contextual mapping to decision-relevant constructs like risk and actionability [5, 6]. Without such interpretation, designers risk misaligning AI insights with practical imperatives, such as ensuring the safety of structural materials or the efficiency of energy storage systems.

Central to this discourse is the distinction between confidence and warrant for action. Confidence, often manifested through calibration metrics or probabilistic scores, represents an AI model's internal assessment of its prediction reliability. For example, a well-calibrated model might assign 80% confidence to a property prediction, implying that 80% of such predictions align with true values in aggregate [3, 7]. Yet, this does not constitute a warrant—a justified basis for proceeding with design actions—because it overlooks broader contextual factors. In materials contexts, warrant encompasses not only statistical reliability but also the implications of errors, domain applicability, and the types of uncertainty. A high-confidence prediction in a familiar domain might warrant prioritization of a material candidate. In contrast, the same confidence level in an extrapolated regime could signal the need to defer or redesign, lest epistemic gaps lead to costly failures [8, 9].

This conflation stems from a historical overemphasis on predictive accuracy in AI development, where uncertainty is treated as a byproduct to be minimized rather than a signal to be leveraged. In materials science, where actions involve resource-intensive syntheses or high-risk deployments, such as in aerospace composites or biomedical implants, interpreting uncertainty becomes imperative for ethical and efficient design [1, 10]. Traditional paradigms, rooted in deterministic engineering, inadequately address AI's probabilistic nature, often leading to either reckless overreliance on models or undue conservatism that stifles innovation [2, 11]. For instance, in high-throughput screening, uninterpreted uncertainty might prompt indiscriminate pursuit of low-confidence candidates, wasting resources, or premature abandonment of promising avenues due to perceived risks.

Moreover, the conceptual separation of uncertainty from confidence and risk highlights actionability as a bridging concept. Actionability assesses whether interpreted uncertainty provides sufficient grounds for specific material actions, such as screening initial candidates from vast libraries, prioritizing those with favorable risk profiles, deferring decisions pending more information, redesigning AI architectures to reduce epistemic components, stopping unviable paths, hedging through parallel explorations, or diversifying material portfolios to mitigate aleatory effects [12, 13]. This view aligns with decision-theoretic principles, in which actions are not merely reactive but are strategically informed by the qualitative and quantitative dimensions of uncertainty.

The need for interpretation is amplified by the interdisciplinary nature of materials AI, which blends computational predictions with physical realizations. Uncertainty in AI outputs—whether from neural networks forecasting band gaps or generative models proposing novel compounds—carries semantic layers, in which meanings shift across design contexts [4, 14]. For example, aleatory uncertainty in stochastic simulations might be tolerable in low-stakes screening but prohibitive in risk-sensitive applications like nuclear materials. Failing to interpret these nuances risks decision failures, such as overconfidence in AI-driven optimizations that ignore tail risks or paralysis in the face of ambiguous signals [15].

This paper synthesizes existing conceptual categories of uncertainty and extends them into a novel decision-theoretic model tailored for materials AI. By framing uncertainty as a design signal, it challenges the prevailing focus on quantification alone and advocates interpretive frameworks that enhance human-AI collaboration [16]. The proposed Uncertainty-to-Action Map offers a structured pathway, incorporating gates to evaluate the implications of uncertainty and avert common pitfalls. Ultimately, this conceptual lens posits that properly interpreted uncertainty not only mitigates risks but also unlocks creative potentials in materials design, fostering resilient and innovative outcomes [17, 18].

This manuscript advances a decision-theoretic interpretation of uncertainty in materials-focused artificial intelligence, rather than a technical framework for uncertainty quantification or model optimization. Its scope is intentionally constrained to AI systems used as decision-support instruments in materials discovery, screening, prioritization, and deployment planning. The framework does not assume autonomous decision-making by AI systems, nor does it prescribe algorithmic implementations, probabilistic estimators, or numerical thresholds.

The proposed Uncertainty-to-Action Map is materials-specific in two critical senses. First, it assumes that AI outputs ultimately inform physical actions—such as synthesis, scale-up, deployment, or rejection—where errors incur asymmetric costs, are irreversible, and have long-term consequences. Second, it presumes heterogeneous design contexts, where

the same uncertainty signal may warrant distinct actions depending on material class, application domain, and consequence severity.

Several non-goals are explicit. This paper does not attempt to:

- Benchmark uncertainty quantification methods or compare Bayesian and non-Bayesian techniques
- Propose new calibration metrics or confidence scores
- Optimize decision policies numerically
- Replace human judgment in materials design decisions

Instead, the contribution is epistemic and interpretive: it clarifies what uncertainty *means* in materials AI, how it differs from confidence and risk, and how it should be translated into justified action or principled restraint. Silence—defined as the deliberate choice not to act—is treated as a valid and sometimes optimal outcome within the framework.

Uncertainty

Uncertainty denotes incomplete knowledge regarding model inputs, internal representations, or predicted outcomes. In materials AI, uncertainty is not a singular quantity but a composite signal arising from multiple sources. Epistemic uncertainty reflects knowledge gaps attributable to limited data, model misspecification, or extrapolation beyond the training domain. Aleatory uncertainty captures irreducible variability inherent to stochastic physical processes or to measurement noise [19-29]. Semantic uncertainty arises when the meaning of an AI output is ambiguous relative to its design intent, as when abstract property predictions fail to cleanly map onto operational requirements.

Confidence

Confidence represents a model's self-assessed reliability with respect to its predictions. It is an internal signal—often probabilistic or score-based—that reflects consistency or calibration under assumed conditions. Importantly, confidence is not action-directive. A confident prediction may still be epistemically fragile, semantically misaligned, or unacceptable under high-stakes conditions [21].

Risk

Risk is defined here as consequence-weighted uncertainty. Unlike uncertainty, which describes an epistemic state, risk evaluates exposure by combining uncertainty with the magnitude, irreversibility, and asymmetry of potential outcomes. In materials contexts, risk cannot be inferred from uncertainty alone [19]; it emerges only when uncertainty is interpreted relative to deployment context, safety margins, environmental impact, and economic cost.

Actionability

Actionability denotes the epistemic warrant to act. It is not synonymous with confidence, low uncertainty, or acceptable risk individually [9]. Rather, actionability arises when interpreted uncertainty supports a specific class of actions—such as screening, prioritization, deferral, redesign, or stopping—given the stakes and objectives. Actionability is therefore contextual, conditional, and graded, not binary. Table 1 summarizes the conceptual distinctions between uncertainty, confidence, risk, and actionability, clarifying their non-equivalence and distinct roles in materials AI decision-making.

Table 1. Conceptual differentiation between uncertainty, confidence, risk, and actionability in materials AI

Concept	Definition (conceptual)	What it is not	Common misinterpretation in materials AI	Decision implication
Uncertainty	Incomplete knowledge about inputs, representations, or outcomes	Error or lack of model quality	Treated as noise to minimize	Signal requiring interpretation
Confidence	Model's self-assessed reliability under assumed conditions	Warrant for action	Equated with safety or correctness	Requires contextual evaluation
Risk	Consequence-weighted uncertainty	Uncertainty magnitude alone	Ignored or inferred directly from confidence	Guides tolerance thresholds
Actionability	Epistemic warrant to act or abstain	Binary decision or automation	Assumed once confidence is high	Contextual, graded, and conditional

Background and Synthesis

Types of uncertainty as conceptual categories

Uncertainty in materials AI manifests in distinct conceptual categories, each with implications for design decisions. Epistemic uncertainty arises from incomplete knowledge, such as limited training data or model approximations, and is theoretically reducible through additional information or refinements [19]. In materials contexts, this might stem from sparse datasets in rare-earth compounds, where gaps in understanding phase behaviors lead to unreliable predictions. Aleatory uncertainty, conversely, reflects inherent randomness, irreducible even with perfect knowledge, as seen in quantum fluctuations or thermal noise affecting property measurements [19]. Semantic uncertainty introduces ambiguity in interpretation, where terms like “stability” or “durability” carry context-dependent meanings, potentially misaligning AI outputs with design intents [20]. These categories interplay; for instance, epistemic elements can amplify aleatory effects in generative design, while semantic layers complicate risk assessments in multifunctional materials [21]. **Table 2** outlines the primary sources, reducibility, and decision implications of epistemic, aleatory, and semantic uncertainty in materials AI.

Table 2. Conceptual uncertainty types in materials AI and their dominant decision implications

Uncertainty Type	Source	Reducibility	Typical Materials Example	Primary Risk Concern	Typical Action Response
Epistemic	Data sparsity, model misspecification, and extrapolation	Reducible	Sparse alloy composition spaces	False confidence	Redesign, defer, targeted data
Aleatory	Inherent variability, stochastic processes	Irreducible	Synthesis variability in nanomaterials	Outcome spread	Hedging, diversification
Semantic	Ambiguity in meaning or intent alignment	Context-dependent	“Stability” vs deployment requirements	Misaligned decisions	Clarification, expert review

Why calibration language can mislead decisions conceptually

Calibration language in AI, emphasizing alignment between predicted probabilities and observed frequencies, can conceptually mislead by fostering a false equivalence between statistical reliability and decision warrant. In materials AI,

well-calibrated models might report high confidence in predictions, yet this overlooks uncertainty's qualitative aspects, such as epistemic origins that signal extrapolation risks [22]. Such language promotes a narrow view, equating calibration with actionability and ignoring how miscalibration in out-of-domain scenarios—common in materials exploration—undermines trust [23]. Conceptually, this misleads by prioritizing aggregate performance over case-specific interpretations, potentially encouraging actions like prioritizing flawed candidates in battery design without addressing underlying ambiguities [24, 26].

Risk is consequence-weighted uncertainty (materials-specific)

Risk conceptually integrates uncertainty with its potential consequences, weighting epistemic, aleatory, and semantic components by their impact on material outcomes. In high-stakes domains such as structural alloys, low-probability epistemic uncertainties can pose catastrophic risks if failure under load is involved [26]. This weighting distinguishes risk from raw uncertainty; for example, aleatory variability in nanomaterial synthesis might pose minimal risk in exploratory screening but elevate to high risk in safety-critical applications [27]. Materials-specific considerations, such as fabrication cost asymmetries or environmental impacts, further modulate this weighting, framing risk as a dynamic posture rather than a static metric [28, 29].

Decision failures: Overconfidence, paralysis, and false humility

Decision failures in materials AI stem from the mishandling of uncertainty signals. Overconfidence occurs when high reported confidence prompts unwarranted actions, disregarding epistemic gaps that could invalidate predictions in novel materials [30]. Paralysis arises from overwhelming uncertainty, halting progress even when actionability exists, as in deferring viable candidates because of unweighted aleatory noise [31]. False humility involves excessive caution, underutilizing reliable AI insights through perpetual redesign or stopping, stemming from semantic misinterpretations that inflate perceived risks [32, 33]. These modes highlight the need for interpretive frameworks to balance boldness and prudence in design.

Proposed framework

The Uncertainty-to-Action Map is introduced as a decision-theoretic conceptual framework that formalizes how uncertainty in materials-focused artificial intelligence should be *interpreted* rather than minimized to support justified design actions. Unlike technical approaches to uncertainty quantification, the framework does not prescribe algorithms, confidence thresholds, or probabilistic estimators. Instead, it provides an epistemic structure that links uncertainty signals to risk postures and, ultimately, to classes of warranted action in materials design.

At its core, the map posits a directional translation pathway. Uncertainty enters the framework not as a scalar output but as a typed signal, explicitly distinguished into epistemic, aleatory, and semantic forms. These uncertainty types do not directly authorize action. Rather, they are first interpreted through an intermediate layer of risk posture formation, where contextual consequences weigh uncertainty. Only after this interpretive step can uncertainty meaningfully inform design actions such as screening large candidate spaces, prioritizing promising materials, deferring judgment pending additional information, redesigning models or representations, halting unproductive pursuits, hedging through parallel strategies, or diversifying portfolios to manage irreducible variability.

The translation from uncertainty to action is governed by a set of conceptual gates that constrain interpretation and prevent premature or unjustified decisions. The *stake assessment gate* evaluates the magnitude, irreversibility, and societal relevance of potential outcomes, separating low-stakes exploratory tasks from high-stakes deployment scenarios. The *ambiguity detection gate* identifies entanglements between semantic and epistemic uncertainty, flagging cases in which unclear meaning, rather than a lack of knowledge, drives apparent risk. The *domain scope gate* assesses whether predictions fall within the model's epistemic enclosure or represent extrapolative claims that demand restraint. The *cost asymmetry gate* evaluates imbalances between upside potential and downside harm, recognizing that identical levels of uncertainty may justify different actions depending on the asymmetry of consequences. Finally, the *stopping logic gate*

establishes conditions under which further uncertainty reduction ceases to be epistemically or practically justified, legitimizing principled abstention.

Through these gates, uncertainty is transformed into qualitative risk postures—such as tolerable, elevated, or prohibitive—which function as decision-relevant summaries rather than numerical scores. These postures then map to action classes in a structured but non-deterministic manner, preserving human judgment while constraining arbitrariness. Importantly, the framework allows for feedback loops: actions such as redesign or deferral can re-enter the map with modified uncertainty profiles, enabling iterative refinement without action drift.

The framework also makes explicit several systematic failure modes that arise when uncertainty is misinterpreted. *Confidence theater* occurs when high reported confidence is mistaken for epistemic sufficiency. *Risk blindness* arises when uncertainty is evaluated without considering consequences. *Semantic neglect* leads to actions misaligned with design intent due to unresolved ambiguity in meaning. *Action drift* arises when decisions shift without reevaluation of interpretive gates. By structurally foregrounding these failure modes, the Uncertainty-to-Action Map positions actionability as a distinct epistemic outcome—earned through interpretation rather than inferred solely from confidence. **Table 3** details the interpretive gates of the Uncertainty-to-Action Map, linking each gate to the failure modes it prevents and the classes of action it enables.

Table 3. Interpretive gates in the uncertainty-to-action map and their role in shaping risk posture and action

Gate	Evaluative question	Primary failure mode prevented	Risk posture impact	Enabled action classes
Stake assessment	What are the consequences of being wrong?	Risk blindness	Elevates or suppresses tolerance	Prioritize, defer, stop
Ambiguity detection	Is uncertainty epistemic or semantic?	Semantic neglect	Reclassifies uncertainty	Clarify, redesign
Domain scope	Is this within epistemic enclosure?	Confidence theater	Flags extrapolation	Defer, abstain
Cost asymmetry	Are the downsides and the upsides balanced?	Overconfidence	Skews posture conservatively	Hedge, diversify
Stopping logic	Is further reduction justified?	Action drift	Terminates loops	Stop, silence

In this way, the framework reframes uncertainty from a liability to be suppressed into a **constructive design signal**, enabling materials AI systems to support decisions that are not only informed, but epistemically warranted and contextually responsible. The model transforms raw ambiguity into a structured decision-making pipeline. As illustrated in **Figure 1**, the framework maps the transition from uncertainty categories to specific action classes through a series of decision nodes.

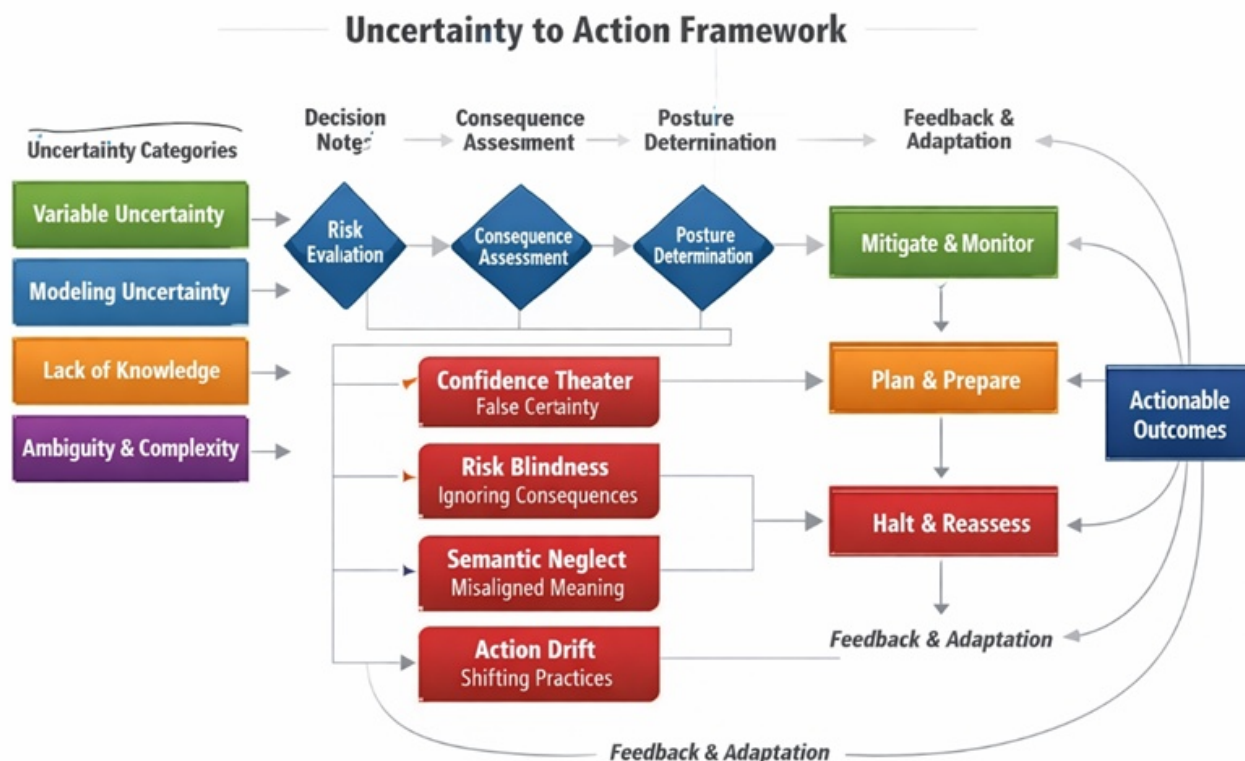


Figure 1. Uncertainty categories entering from the left, passing through gates to yield risk postures, which map rightward to action classes, with feedback loops for iterative refinement

Results and Discussion

The Uncertainty-to-Action Map presented in this manuscript offers a conceptual lens for navigating the complexities of uncertainty in materials AI, transforming it from a potential liability into a strategic design signal. This framework underscores the importance of distinguishing uncertainty, confidence, risk, and actionability, thereby addressing gaps in traditional approaches that often prioritize predictive accuracy at the expense of interpretive depth [1, 2]. By incorporating gates such as stake assessment and ambiguity detection, the map facilitates a more nuanced decision process, enabling materials scientists to align AI outputs with real-world design imperatives. For instance, in scenarios involving high-stakes materials such as those used in nuclear applications, the asymmetry of costs gate highlights how minor epistemic uncertainties can escalate into significant risks, prompting actions such as hedging or diversification rather than outright prioritization [3, 4].

One key implication of this model is its potential to mitigate common decision failures in materials AI workflows. Overconfidence, often exacerbated by misleading calibration language, can be countered through the domain scope gate, which evaluates whether predictions fall within well-characterized regimes [5, 19]. Similarly, paralysis from unweighted uncertainties is alleviated by the risk posture evaluation, which consequence-weights aleatory and epistemic components to discern tolerable from prohibitive risks [6, 21]. False humility, where designers err on the side of excessive caution, is addressed by the stopping logic gate, which provides clear criteria for when uncertainty-reduction efforts yield diminishing returns [7, 8]. These mechanisms collectively promote a balanced approach, fostering innovation while safeguarding against undue exposure in materials discovery pipelines.

However, the framework's conceptual nature invites consideration of its limitations. As a decision-theoretic model, it assumes access to categorized uncertainty types, yet in practice, disentangling epistemic, aleatory, and semantic

uncertainties may require supplementary tools not detailed here [9, 10]. Furthermore, the map's gates rely on human judgment for parameterization, potentially introducing subjective biases that could undermine its objectivity [11, 12]. In interdisciplinary settings, where materials AI intersects with domains like chemistry or physics, semantic uncertainties may be particularly prone to misinterpretation, as terms like "reliability" vary across contexts [13, 14]. These limitations highlight the need for complementary empirical validations, though this manuscript remains focused on theoretical foundations.

Future directions for extending this framework are abundant. Integrating it with emerging AI paradigms, such as federated learning for distributed materials datasets, could enhance its applicability to collaborative design environments [15, 16]. Additionally, exploring synergies with explainable AI techniques might refine the ambiguity detection gate, offering insights into why certain uncertainties arise [17, 18]. Conceptually, the model could be adapted to other uncertainty-laden fields, such as drug discovery or climate modeling, where similar action classes apply [21, 22]. Ultimately, by emphasizing uncertainty as a design asset, this framework paves the way for more resilient AI-assisted materials innovation, encouraging a shift from minimizing to strategically leveraging incomplete knowledge [23, 24].

Conclusion

This conceptual manuscript has elucidated the critical role of interpreting uncertainty as a design signal in materials AI, distinguishing it from confidence, risk, and actionability to inform judicious design actions. Traditional conflations of these concepts have led to suboptimal decisions. Still, the proposed Uncertainty-to-Action Map provides a decision-theoretic pathway that incorporates gates to translate uncertainty types into risk postures and action classes such as screening, prioritizing, or stopping. By mitigating failure modes such as confidence theater and risk blindness, the framework fosters robust, context-aware decision-making in materials science.

In summary, viewing uncertainty through this lens not only reduces vulnerabilities but also unlocks opportunities for creative exploration, ensuring AI serves as a reliable partner in advancing materials innovation. Future conceptual refinements could further enhance its utility across scientific domains.

Acknowledgements

None

Conflict of interest

None

Financial support

None

Ethics statement

None

Received: 12 Aug 2023

Revised: 11 Nov 2023

Accepted: 03 Dec 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

- Korolev V, Nevolin I, Protsenko P. A universal similarity based approach for predictive uncertainty quantification in materials science. *Sci Rep.* 2022;12:14931.
- Wen M, Tadmor EB. Uncertainty quantification in molecular simulations with dropout neural network potentials. *npj Comput Mater.* 2020;6:124.
- Zhang H, Chen W, Iyer A, Apley DW, Chen W. Uncertainty-aware mixed-variable machine learning for materials design. *Sci Rep.* 2022;12:19760.
- Tan AR, Urata S, Goldman S, Dietschreit JCB, Gómez-Bombarelli R. Single-model uncertainty quantification in neural network potentials does not always mean the same. *npj Comput Mater.* 2023;9:225.
- Tavazza F, DeCost B, Choudhary K. Uncertainty Prediction for Machine Learning Models of Material Properties. *ACS Omega.* 2021;6:32431-40.
- Janet JP, Duan C, Yang T, Nandy A, Kulik HJ. Leveraging uncertainty from deep learning for trustworthy materials discovery. *ACS Omega.* 2021;6:33149-58.
- Scalia G, Grambow CA, Pernici B, Li YP, Green WH. Uncertainty quantification using neural networks for molecular property prediction. *J Chem Inf Model.* 2020;60:3775-85.
- Heid E, McGill CJ, Vermeire FH, Green WH. Characterizing uncertainty in machine learning for chemistry. *J Chem Inf Model.* 2023;63:4012-29.
- Yin T, Panapitiya G, Coda ED, Saldanha EG. Evaluating uncertainty-based active learning for accelerating the generalization of molecular property prediction. *J Cheminform.* 2023;15:105.
- Chen H, Frey NC, Bowman D, et al. Machine learning-based inverse design for electrochemically active molecules. *Proc Natl Acad Sci U S A.* 2022;119(32):e2206321119.
- Zhou G, Lubbers N, Barros K, Tretiak S, Nebgen B. Deep learning of dynamically responsive chemical Hamiltonians with semiempirical quantum mechanics. *Proc Natl Acad Sci U S A.* 2022;119(27):e2120333119.
- Mansbach RA, Ferguson AL, Kilian KA, Keten S, Lequieu J, de Pablo JJ, et al. Conformal prediction under feedback covariate shift for machine learning on streaming data. *Proc Natl Acad Sci U S A.* 2022;119(28):e2204569119.
- Blei DM, Kucukelbir A, McAuliffe JD. Comparing methods for statistical inference with model uncertainty. *Proc Natl Acad Sci U S A.* 2021;118(10):e2120737118.
- Han H, Wang W-Y, Mao B-H. Automated crystal system identification from electron diffraction patterns using multimodal data fusion. *Proc Natl Acad Sci U S A.* 2023;120(45):e2309240120.

Angelino E, Larus-Stone N, Alabi D, Seltzer M, Rudin C. Cross-prediction-powered inference. *Proc Natl Acad Sci U S A*. 2023;120(38):e2322083120.

Tavazza F, Choudhary K, DeCost B. Approaches for Uncertainty Quantification of AI-predicted Material Properties: A Comparison. *arXiv preprint arXiv:2310.13136*. 2023.

Basu K, Hao J, Hintz D, Shah D, Palmer A, Hora GS, et al. Uncertainty quantification methods for ML-based surrogate models of scientific applications. *NeurIPS Workshop*. 2022.

Otis R. Uncertainty reduction and quantification in computational thermodynamics. *Comput Mater Sci*. 2022;212:111590.

Hüllermeier E, Waegeman W. Aleatoric and epistemic uncertainty in machine learning: An introduction to concepts and methods. *Mach Learn*. 2021;110:457-506.

Nemani V, Biggio L, Huan X, Hu Z, Fink O, Tran A, et al. Uncertainty quantification in machine learning for engineering design and health prognostics: A tutorial. *Mech Syst Signal Process*. 2023;205:110796.

Wimmer L, Sale Y, Hofman P, Bischl B, Hüllermeier E. Quantifying Aleatoric and Epistemic Uncertainty in Machine Learning: Are Conditional Entropy and Mutual Information Appropriate Measures? *PMLR*. 2023;216:2282-92.

Boström H. Aleatoric and epistemic uncertainty with random forests. *Adv Intell Data Anal XVIII*. 2020:444-56.

National Institute of Standards and Technology. Artificial Intelligence Risk Management Framework (AI RMF 1.0). NIST. 2023.

Aristi Baquero J, Burkhardt R, Govindarajan A, Wallace T. Derisking AI by design: How to build risk management into AI development. *McKinsey*. 2020.

Hariri-Ardebili MA, Salazar F. Engaging soft computing in material and modeling uncertainty quantification of dam engineering problems. *Soft Comput*. 2020;24:11583-604.

Chu K, Foster ME, Sills RB, Zhou X, Zhu T, McDowell DL. Temperature and composition dependent screw dislocation mobility in austenitic stainless steels from large-scale molecular dynamics. *npj Comput Mater*. 2020;6:179.

Miller BK, Geiger M, Smidt TE, Noé F. Relevance of rotationally equivariant convolutions for predicting molecular properties. *arXiv preprint arXiv:2008.08461*. 2020.

Yao S, Halpern Y, Thain N, Wang X, Lee K, Prost F, et al. Measuring Recommender System Effects with Simulated Users. *arXiv preprint arXiv:2101.04526*. 2021.

Bellerive A, Kaminski J, Lewis PM, Colas P, Diener R, Peter Kluit P, et al. MPGDs for TPCs at future lepton colliders. *arXiv preprint arXiv:2203.06267*. 2022.

Malik B, Kashyap AR, Kan M-Y, Poria S. UDAPTER --Efficient Domain Adaptation Using Adapters. *arXiv preprint arXiv:2302.03194*. 2023.

Bechetti DH, Semple JK, Zhang W, Fisher CR. Temperature-Dependent Material Property Databases for Marine Steels—Part 1: DH36. *Integr Mater Manuf Innov*. 2020;9:257-86.

Janet JP, Kulik HJ. Resolving transition metal chemical space: feature selection for machine learning and structure-property relationships. *J Phys Chem A*. 2020;124:1850-61.

Pernot P, Cailliez F. A critical review of statistical calibration methods for estimating thermochemical properties of materials. *J Phys Chem A*. 2021;125:6245-58.