
Surrogate Modeling for the Design of Optimal Lattice Structures using Tensor Completion

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 When designing new materials, it is often necessary to design a material with
2 specific desired properties. Unfortunately, as new design variables are added, the
3 search space grows exponentially, which makes synthesizing and validating the
4 properties of each material very impractical and time-consuming. In this work,
5 we focus on the design of optimal lattice structures with regard to mechanical
6 performance. Computational approaches, including the use of machine learning
7 (ML) methods, have shown improved success in accelerating materials design.
8 However, these ML methods are still lacking in scenarios when training data (i.e.
9 experimentally validated materials) come from a non-uniformly random sampling
10 across the design space. For example, an experimentalist might synthesize and vali-
11 date certain materials more frequently because of convenience. For this reason, we
12 suggest the use of tensor completion as a surrogate model to accelerate the design
13 of materials in these atypical supervised learning scenarios. In our experiments, we
14 show that tensor completion is superior to classic ML methods such as Gaussian
15 Process and XGBoost with biased sampling of the search space, with around 5%
16 increased R^2 . Furthermore, tensor completion still gives comparable performance
17 with a uniformly random sampling of the entire search space.

18 1 Introduction

19 Efficiently designing new materials with optimal properties can be a significant challenge due to
20 the explosive number of combinations as new design variables are introduced. When searching
21 for new materials with particular material property values, it can be very expensive to exhaustively
22 generate these combinations of materials. For this reason, there is growing interest in being able to
23 predict these material property values, without having to first produce the material. Computational
24 methods have shown great promise for inferring material property values, allowing material scientists
25 to quickly find materials with optimal properties. This includes the use of Density Functional Theory
26 (DFT) [8, 22, 19, 15], machine learning (ML) [3, 6, 9, 10, 13, 16, 23, 25], or both [20]. In this work,
27 we focus on designing optimal lattice structures, with regards to mechanical performance [6, 7].

28 Unfortunately, these ML methods still underachieve in scenarios where the training data does not
29 come from a uniformly random sample of the entire search space. This scenario may arise in
30 practice when an experimentalist might synthesize and validate certain materials of the design space
31 more than others, due to convenience. In order to overcome this, we abstract away the various
32 components of material design and instead model them as instances of tensor completion. Here,
33 each design parameter to be optimized corresponds to an individual mode of our tensor, so we can
34 leverage tensor completion algorithms to predict the entirety of the search space. In our experiments,
35 tensor completion methods provide nearly state-of-the-art performance in typical supervised learning
36 scenarios. Furthermore, we simulate an experimentalist performing experiments to serve as training

data by unevenly sampling across design variables. In these scenarios, tensor completion gives better results than other baseline ML methods, including Gaussian Process [24], XGBoost [5], and some ensembles of multiple methods. Of these baseline methods, Gaussian Process performed the best in our experiments so we show results compared against this.

2 Methods

In this section we give a detailed description on the methods used in this work, including our pipeline for surrogate modeling with tensor completion, its training, and the dataset used for evaluation.

2.1 Tensor Completion for Surrogate Modeling

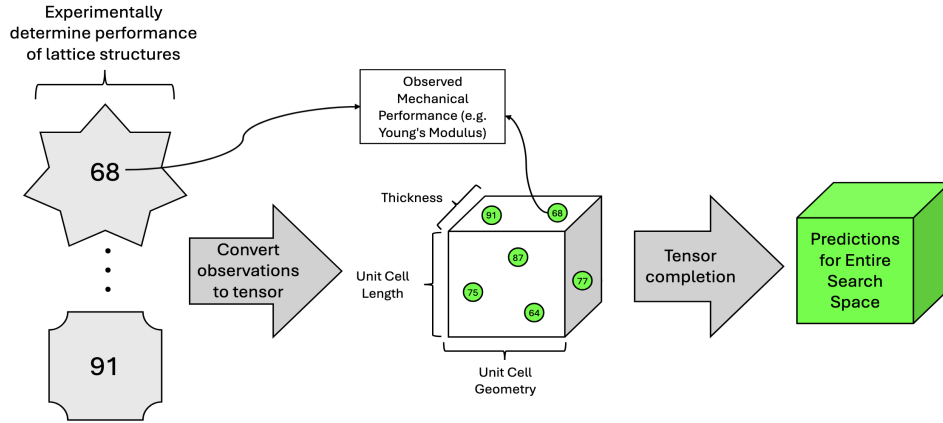


Figure 1: We present various material design problems as instances of tensor completion, in order to abstract away the various design components and use tensor methods to infer the entire search space.

In Figure 1, we show the complete surrogate modeling pipeline we are proposing. After experimentally determining the performance for various material structures, we can convert those observations to entries of a tensor and use tensor completion to infer the performance of the remainder of the design search space.

2.1.1 Tensors & Tensor Decomposition

Tensors are a general term for multidimensional arrays. In other words, a 1st order tensor is just a vector, and a 2nd order tensor is a matrix. In this work, we will be looking at higher order tensors, denoted $\mathcal{X}$. Tensor decomposition is a general term for expressing a tensor as several smaller factors, and is an extension of matrix factorization to multidimensional datasets. A common form of tensor decomposition is the Canonical Polyadic Decomposition (CPD) [12, 21]. CPD expresses a tensor as a sum of rank-one tensors. A rank R CPD decomposition of a third-order tensor $\mathcal{X} \in \mathbb{R}^{I \times J \times K}$ would be expressed as: $\mathcal{X} \approx \sum_{r=1}^R (\mathbf{a}_r \circ \mathbf{b}_r \circ \mathbf{c}_r)$, where $\circ$ denotes outer product, $\mathbf{a}_r \in \mathbb{R}^I$, $\mathbf{b}_r \in \mathbb{R}^J$, and $\mathbf{c}_r \in \mathbb{R}^K$.

Tensor decomposition is a powerful tool, not only for data compression but also for analyzing multidimensional datasets. For example, the smaller tensor factors are commonly used for pattern discovery in these multidimensional datasets, and further more for tensor completion. Tensor completion is the process of filling in the missing values of a tensor.

2.1.2 Tensor Completion Training

We illustrate the training of our tensor completion models in Figure 2. We first randomly initialize the factor matrices to generate an initial (random) reconstruction of the full tensor. Then we use error of the observed tensor values and the associated predictions to iteratively get a more accurate reconstruction of the full tensor. For our error function, we use Mean Absolute Error (MAE) =

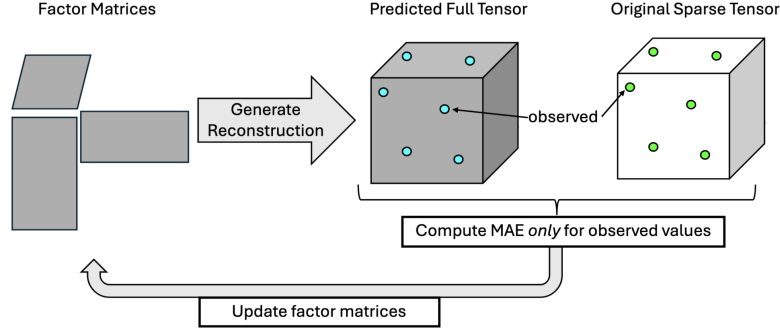


Figure 2: Here we visualize the training process, where we iteratively update our factor matrices to produce a better tensor reconstruction. The gray regions in the factor matrices and predicted tensor correspond to predicted values for which we do not have the ground truth during training.

67 $\frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i|$. To optimize the parameters (factor matrices) with regards to this error term, we
 68 use the Adam optimizer [11].

69 2.1.3 Tensor Completion Models

70 There are several classes of tensor completion methods. There are classical methods such as CPD
 71 [12] and Tucker [4], and there are more advanced neural tensor completion methods such as NeAT
 72 [2] or CoSTCo [14], which leverage the use of neural networks in addition to tensor decomposition.

73 For our experiments, we use CPD-S (a CPD-based method that imposes smoothness constraints on
 74 some of the tensor factors) [1, 17], NeAT [2], and an ensemble tensor completion method to aggregate
 75 the results of multiple tensor completion algorithms. For our results, we only display an ensemble of
 76 CPD-S with different ranks, and a NeAT model as it showed the best performance on this dataset.

77 2.2 Baseline Comparisons

78 In this work, we compare our tensor completion methods against Gaussian Process, XGBoost, Multi-
 79 Layer Perceptrons (MLPs), and using an ensemble of several ML methods. Due to limited space, we
 80 only show the results for Gaussian Process, as it showed the best results of the baseline comparisons
 81 on this dataset.

82 2.3 Dataset Description

83 To verify the utility of using tensor completion as a surrogate model for lattice structure design, we
 84 use the dataset from Gongora et al. [6] with various lattice structures’ design components and the
 85 corresponding mechanical performance. Mechanical performance is characterized using the Young’s
 86 modulus (E) values and the ratio between E and the design’s mass m ($\tilde{E}$). These are essential
 87 mechanical properties to describe the elastic behavior of a structure under an applied force [6].

88 3 Experiments

89 In this section, we experimentally validate our method in a variety of settings, with uniformly random
 90 and biased random sampling of the trainset, with varying training sizes.

91 3.1 Tensor Completion Performance

92 From Figure 3 we see that tensor completion gives good results in typical supervised learning tasks
 93 (i.e. uniform random sampling). However, in this scenario, we can achieve better results with less
 94 data using a Gaussian Process.

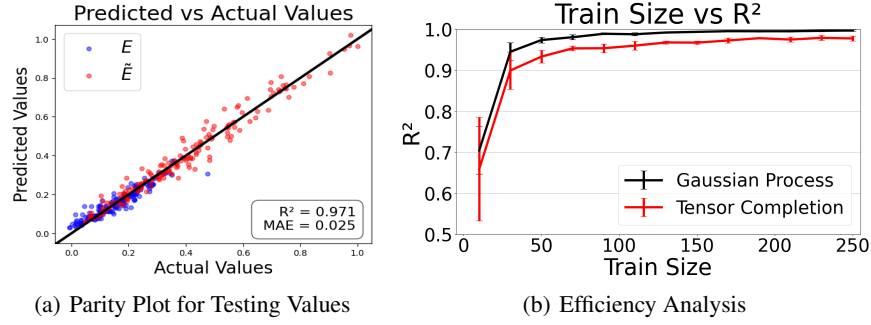


Figure 3: In (a) we display a parity plot for predicting E and $\tilde{E}$ simultaneously. We use 100 train and 170 test values. In (b) we display the R^2 for a uniform random sampling of different training sizes. We show the mean and standard deviation on the test set, using 5 iterations for each train size.

3.2 Biased Sampling

Here we observe the performance of surrogate modeling when the training data is not uniformly randomly sampled throughout the search space. We sample a different number of training values for each value of a design variable. The goal here is to simulate an experimentalist who might have conducted more experiments in a specific region of the search space, rather than uniformly sampling throughout the entire search space. This is useful if a certain region of the search space is cheaper or more convenient (e.g. with regards to time or money) to experimentally validate than others.

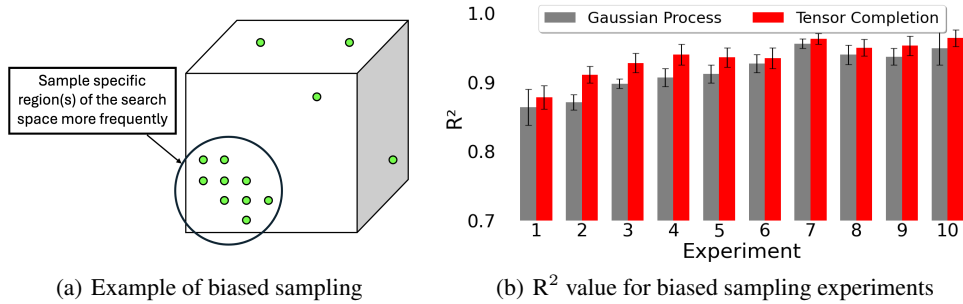


Figure 4: In (a) we visualize how the search space might be sampled in a biased manner. In (b), we conduct experiments for a biased-random sampling of the search space. We show the average and standard deviation of the R^2 for 5 iterations. Going from experiment 1 to 10, we decrease the bias in sampling by making the range smaller. Details on the generation of number of training samples and exact number of training samples for each value are discussed in the appendix section in Table 1.

From Figure 4 we observe that tensor completion methods are better in handling a biased sampling in the training set, and less prone to overfitting to the observed values. This makes tensor completion very useful for a wider variety of real-world scenarios than standard supervised learning techniques.

4 Conclusion

In this work, we cast the design of optimal lattice structures with respect to mechanical performance as an example of a tensor completion problem. In this way, we are able to handle both typical (uniformly random sampling) and atypical (biased sampling) supervised learning problems in order to accelerate the design of optimal lattice structures. For these reasons, tensor completion could be a promising surrogate model to accelerate the search for optimal lattice structures with exponentially increasing design space.

References

- [1] Dawon Ahn, Jun-Gi Jang, and U Kang. Time-aware tensor decomposition for sparse tensors. In *2021 IEEE 8th International Conference on Data Science and Advanced Analytics (DSAA)*, pages 1–2. IEEE, 2021.
- [2] Dawon Ahn, Uday Singh Saini, Evangelos E Papalexakis, and Ali Payani. Neural additive tensor decomposition for sparse tensors. In *33rd ACM International Conference on Information and Knowledge Management*. ACM, 2024.
- [3] Kirstin Alberi, Marco Buongiorno Nardelli, Andriy Zakutayev, Lubos Mitas, Stefano Curtarolo, Anubhav Jain, Marco Fornari, Nicola Marzari, Ichiro Takeuchi, Martin L Green, et al. The 2019 materials by design roadmap. *Journal of Physics D: Applied Physics*, 52(1):013001, 2018.
- [4] Ivana Balažević, Carl Allen, and Timothy M Hospedales. Tucker: Tensor factorization for knowledge graph completion. In *Empirical Methods in Natural Language Processing*, 2019.
- [5] Tianqi Chen and Carlos Guestrin. Xgboost: A scalable tree boosting system. In *Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining*, pages 785–794, 2016.
- [6] Aldair E Gongora, Caleb Friedman, Deirdre K Newton, Timothy D Yee, Zachary Doorenbos, Brian Giera, Eric B Duoss, Thomas Y-J Han, Kyle Sullivan, and Jennifer N Rodriguez. Accelerating the design of lattice structures using machine learning. *Scientific Reports*, 14(1):13703, 2024.
- [7] Aldair E Gongora, Siddharth Mysore, Beichen Li, Wan Shou, Wojciech Matusik, Elise F Morgan, Keith A Brown, and Emily Whiting. Designing composites with target effective young’s modulus using reinforcement learning. In *Proceedings of the 6th annual ACM symposium on computational fabrication*, pages 1–11, 2021.
- [8] Jürgen Hafner, Christopher Wolverton, and Gerbrand Ceder. Toward computational materials design: the impact of density functional theory on materials research. *MRS bulletin*, 31(9):659–668, 2006.
- [9] Yousof Haghshenas, Wei Ping Wong, Vidhyasaharan Sethu, Rose Amal, Priyank Vijaya Kumar, and Wey Yang Teoh. Full prediction of band potentials in semiconductor materials. *Materials Today Physics*, 46:101519, 2024.
- [10] Jeffrey Hu, David Liu, Nihang Fu, and Rongzhi Dong. Realistic material property prediction using domain adaptation based machine learning. *Digital Discovery*, 3(2):300–312, 2024.
- [11] Diederik P Kingma. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- [12] T.G. Kolda and B.W. Bader. Tensor decompositions and applications. *SIAM review*, 51(3), 2009.
- [13] Qin Li, Nihang Fu, Sadman Sadeed Omee, and Jianjun Hu. Md-hit: Machine learning for material property prediction with dataset redundancy control. *npj Computational Materials*, 10(1):245, 2024.
- [14] Hanpeng Liu, Yaguang Li, Michael Tsang, and Yan Liu. Costco: A neural tensor completion model for sparse tensors. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, pages 324–334, 2019.
- [15] Priyanka Makkar and Narendra Nath Ghosh. A review on the use of dft for the prediction of the properties of nanomaterials. *RSC advances*, 11(45):27897–27924, 2021.
- [16] Shaan Pakala, Dawon Ahn, and Evangelos Papalexakis. Tensor completion for surrogate modeling of material property prediction. *arXiv preprint arXiv:2501.18137*, 2025.

- [17] Shaan Pakala, Bryce Graw, Dawon Ahn, Tam Dinh, Mehnaz Tabassum Mahin, Vassilis Tsotras, Jia Chen, and Evangelos E. Papalexakis. Automating data science pipelines with tensor completion. In *2024 IEEE International Conference on Big Data (BigData)*, pages 1075–1084, 2024.
- [18] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot, and E. Duchesnay. Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research*, 12:2825–2830, 2011.
- [19] Gabriel R Schleder, Antonio CM Padilha, Carlos Mera Acosta, Marcio Costa, and Adalberto Fazzio. From dft to machine learning: recent approaches to materials science—a review. *Journal of Physics: Materials*, 2(3):032001, 2019.
- [20] Atsuto Seko, Tomoya Maekawa, Koji Tsuda, and Isao Tanaka. Machine learning with systematic density-functional theory calculations: Application to melting temperatures of single-and binary-component solids. *Physical Review B*, 89(5):054303, 2014.
- [21] Nicholas D Sidiropoulos, Lieven De Lathauwer, Xiao Fu, Kejun Huang, Evangelos E Papalexakis, and Christos Faloutsos. Tensor decomposition for signal processing and machine learning. *IEEE Signal Processing Magazine*, 2016.
- [22] Pragya Verma and Donald G Truhlar. Status and challenges of density functional theory. *Trends in Chemistry*, 2(4):302–318, 2020.
- [23] Yuhao Wang, Yefan Tian, Tanner Kirk, Omar Laris, Joseph H Ross Jr, Ronald D Noebe, Vladimir Keylin, and Raymundo Arróyave. Accelerated design of fe-based soft magnetic materials using machine learning and stochastic optimization. *Acta Materialia*, 194:144–155, 2020.
- [24] Christopher Williams and Carl Rasmussen. Gaussian processes for regression. *Advances in neural information processing systems*, 8, 1995.
- [25] Ya Zhuo, Aria Mansouri Tehrani, and Jakoah Brgoch. Predicting the band gaps of inorganic solids by machine learning. *The journal of physical chemistry letters*, 9(7):1668–1673, 2018.

A Appendix

A.1 Limitations

The main limitation of this work is considering typical supervised learning scenarios, when the training data comes from a uniformly random sampling of the entire data or search space. In our experiments, we see that a Gaussian Process exhibits superior performance in these scenarios.

A.2 Experimental Details

All code used for experiments can be found at <https://anonymous.4open.science/r/Surrogate-Modeling-for-the-Design-of-Optimal-Lattice-Structures-using-Tensor-Completion-36BF>. However, we cannot release the dataset due to the dataset owner’s institutional regulations. More dataset details can be found on the original paper [6].

A.2.1 Hardware

Experiments were fairly lightweight to run. They all were conducted using a MacBook Air 2020 laptop with an Apple M1 chip (8GB).

A.2.2 ML Models

For our experiments, we used an ensemble tensor completion method consisting of a NeAT Rank 24 & 32, CPD Ranks 1, 2, & 4, and CPD-S Ranks 1, 2, & 4. Each model was trained using Adam optimizer [11] with a learning rate and weight decay of 0.01, using MAE as the objective function. A random forest regression using 100 trees was used to aggregate the results of the tensor models.

For our baseline methods, we tried a variety of XGBoost, MLPs, and Gaussian Process models using Scikit-Learn [18]. The best of the baselines (and only one shown in the experiments due to space constraints) was a Gaussian Process with kernel = ConstantKernel(1.0, (1e-3, 1e3)) $\times$ RBF(1.0, (1e-2, 1e2)) + WhiteKernel(1e-3, (1e-5, 1e1)) using Scikit-Learn’s Gaussian Process kernel functions, and an alpha value = 0.01. These were determined empirically by comparing the performance of various hyperparameter values.

A.2.3 Section 3.2 Biased Sampling details

Experiment	Gyroid	Schwarz	Diamond	Lidinoïd	Split P	Range
1	40	21	7	6	3	[3, 40]
2	14	21	5	5	40	[5, 40]
3	8	23	7	15	40	[7, 40]
4	9	24	12	19	40	[9, 40]
5	11	21	16	40	25	[11, 40]
6	14	13	40	22	35	[13, 40]
7	36	15	28	40	16	[15, 40]
8	40	17	21	28	37	[17, 40]
9	30	35	29	19	40	[19, 40]
10	34	40	32	39	21	[21, 40]

Table 1: Number of training samples used for each experiment in Figure 4. Each unit cell geometry type (i.e. tensor slice) has 54 total values (54 values $\times$ 5 slices = 270 total values in search space). For each experiment, the number of training samples for each unit cell geometry type comes from an exponential distribution with scale = 1. This was to create a high disparity in the number of samples for each value. We then converted this exponential distribution to be in the range $[l, 40]$, where $l = 3 + 2 \times e_{num}$ (e_{num} is the current experiment number), and then rounded to the nearest integer value. This was to iteratively decrease the bias in sampling, by making the range of values smaller (i.e. increasing the lower bound l) as we go from experiment 1 to experiment 10. We start with experiment 1 from an exponential distribution on the range $[3, 40]$, and we end with experiment 10 on the range $[21, 40]$.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope?

Answer: [\[Yes\]](#)

Justification: We state that our tensor methods are nearly (not better than) state-of-the-art in standard supervised learning scenarios, where we train on a uniformly random sample of the entire data. We do suggest that tensor methods could offer improvement in biased sampling scenarios, which may represent an actual experimentalists’ training data, and we conduct experiments in Section 3.2 to demonstrate this.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We are transparent in the fact that standard supervised learning models (e.g. Gaussian Process) may outperform tensor models in “regular” supervised learning problems, with a uniform random sampling of the search space.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren’t acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA]

Justification: We do not have any theoretical results, we just discuss empirical results which we conduct experiments to justify.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [No]

Justification: Justification: We give all code required to reproduce the results; however, due to the original dataset owner's institutional regulations, we are not allowed to release the dataset. Further details on the dataset can be found in the original paper [6].

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).

- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [No]

Justification: Justification: We give all code required to reproduce the results; however, due to the original dataset owner’s institutional regulations, we are not allowed to release the dataset. Further details on the dataset can be found in the original paper [6].

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: We discuss all training details for the experiments (further details in the appendix section).

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: For each experiment (and sub-experiment), we perform 5 iterations and display both the mean and error bars (standard deviation) of the 5 iterations.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We cover this in the appendix section.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We have reviewed the code of ethics. We do not use any human research subjects, provide any code/software that could inflict any harm, and we are as transparent as possible with the experiments and results in this work.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: We discuss how using tensor methods to navigate the search space may be beneficial to experimentalists who conduct experiments in a biased manner. This may limit the required costs of generating a uniformly random sample of the entire search space and allow a more convenient navigation of the search space.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: This paper poses no such risk, we are just suggesting a tool that may be useful for experimentalists to accelerate the search of an exponentially large search space for developing optimal materials.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We use the dataset from the work of Gongora et al. [6] and discuss this in the dataset description. Besides this, the code used was either developed by the authors of this work, or properly cited.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: No new assets are released in this work.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: No crowdsourcing or research with human subjects was conducted in this work.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

521 Question: Does the paper describe potential risks incurred by study participants, whether
522 such risks were disclosed to the subjects, and whether Institutional Review Board (IRB)
523 approvals (or an equivalent approval/review based on the requirements of your country or
524 institution) were obtained?

525 Answer: [NA]

526 Justification: No research with human subjects was conducted in this work.

527 Guidelines:

- 528 • The answer NA means that the paper does not involve crowdsourcing nor research with
529 human subjects.
- 530 • Depending on the country in which research is conducted, IRB approval (or equivalent)
531 may be required for any human subjects research. If you obtained IRB approval, you
532 should clearly state this in the paper.
- 533 • We recognize that the procedures for this may vary significantly between institutions
534 and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the
535 guidelines for their institution.
- 536 • For initial submissions, do not include any information that would break anonymity (if
537 applicable), such as the institution conducting the review.

538 16. Declaration of LLM usage

539 Question: Does the paper describe the usage of LLMs if it is an important, original, or
540 non-standard component of the core methods in this research? Note that if the LLM is used
541 only for writing, editing, or formatting purposes and does not impact the core methodology,
542 scientific rigor, or originality of the research, declaration is not required.

543 Answer: [NA]

544 Justification: LLMs were not used to develop any core methods.

545 Guidelines:

- 546 • The answer NA means that the core method development in this research does not
547 involve LLMs as any important, original, or non-standard components.
- 548 • Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>)
549 for what should or should not be described.