

Active-learning and materials design: the example of high glass transition temperature polymers

Chiho Kim[†], Anand Chandrasekaran[†], Anurag Jha, and Rampi Ramprasad, School of Materials Science and Engineering, Georgia Institute of Technology, 771 Ferst Drive NW, Atlanta, GA 30332, USA

Address all correspondence to Rampi Ramprasad at rampi.ramprasad@mse.gatech.edu

(Received 16 January 2019; accepted 29 May 2019)

Abstract

Machine-learning (ML) approaches have proven to be of great utility in modern materials innovation pipelines. Generally, ML models are trained on predetermined past data and then used to make predictions for new test cases. Active-learning, however, is a paradigm in which ML models can direct the learning process itself through providing dynamic suggestions/queries for the “next-best experiment.” In this work, the authors demonstrate how an active-learning framework can aid in the discovery of polymers possessing high glass transition temperatures (T_g). Starting from an initial small dataset of polymer T_g measurements, the authors use Gaussian process regression in conjunction with an active-learning framework to iteratively add T_g measurements of candidate polymers to the training dataset. The active-learning framework employs one of three decision making strategies (exploitation, exploration, or balanced exploitation/exploration) for selection of the “next-best experiment.” The active-learning workflow terminates once 10 polymers possessing a T_g greater than a certain threshold temperature are selected. The authors statistically benchmark the performance of the aforementioned three strategies (against a random selection approach) with respect to the discovery of high- T_g polymers for this particular demonstrative materials design challenge.

Introduction

In order to design new materials for specific applications, we often have to search for materials which possess a given set of properties within a required window. For example, design of polymers for energy storage applications requires that such materials possess simultaneously high dielectric constant and bandgap.^[1–9] Another example is the design of solid polymer electrolytes for Li-ion batteries. Such materials are required to possess a suitably high Li-ion conductivity^[10] in conjunction with an appropriate electrochemical window.^[11] In order to guide this search for materials with promising functionalities, scientists and researchers often rely on intuition gained from past experiments to design the next set of experiments. Even so, how does one decide whether to continue searching within a particularly promising class of materials or switch to searching for candidates in a more unexplored region of chemical or structural space? Rather than making such decisions based purely on human intuition, active-learning algorithms that exploit Bayesian optimization (BO) frameworks may be utilized.^[12–15]

Over the past decade, machine-learning (ML)-based algorithms and techniques have been of tremendous utility in a variety of fields, including in materials science.^[16–19] Approaches to the ML can roughly be divided into two categories, passive and active, each making characteristic assumptions about the

learner and its environment.^[20] In passive learning approaches such as classification, clustering, and regression, the ML algorithm or surrogate model can only make inferences on the environment based on the predetermined training data provided to it. Within the active-learning paradigm, however, the learning algorithm can make dynamic queries or suggestions to direct the learning process itself. For instance, when faced with a paucity of training data, the ML algorithm can direct the user to provide data from unexplored regions (quantified by the prediction) so as to gain knowledge and improve the overall accuracy of the model. On the other hand, provided with some prior user-defined objective/cost function, the algorithm can also direct the user to sample points which can maximize/minimize this function. The former approach of sampling unexplored regions with large uncertainty is referred to as “exploration” whereas the latter approach of acquiring data which maximizes the task-specific utility of a particular action/measurement is referred to as “exploitation.” Another common strategy is to balance exploration and exploitation, by taking both knowledge gain and the task-specific utility of actions into account.

In the current report, we outline an active-learning approach to efficiently search the chemical space for polymers which possess a glass transition temperature (T_g) higher than a certain minimum threshold. In other words, given an initial (small) dataset of polymers and their corresponding T_g measurements, our algorithm provides dynamic suggestions on the next set of polymers to synthesize and test so as to achieve our objective of

[†] Chiho Kim and Anand Chandrasekaran equally contributed to this work.

designing high- T_g polymers. We demonstrate this approach on a dataset of 736 polymer T_g measurements that we have curated from various publicly available sources of data.^[21–23] Our results indicate that all three strategies (exploitation, exploration, and balance exploitation/exploration) consistently outperform a random search approach. We observe that the two best-performing strategies are the balanced exploration/exploitation and pure exploitation strategies. We notice that the former is the most robust strategy to employ in finding the most number of high- T_g polymers using the least number of “experiments,” particularly when the size of the initial dataset is small. The active-learning framework and strategies detailed in this work can be generalized to many materials classes and can also be used to optimize more than one materials property simultaneously.

Method

First, we randomly select five polymers from the dataset of 736 polymers, and pretend that we know the T_g value only for these polymers. We then train a Gaussian process regression (GPR) model on just the five randomly selected polymers; the experimental T_g measurements for the remaining 731 polymers are kept hidden to the learning algorithm. Although this initial model is likely highly inaccurate, it provides us with the ability to make predictions of T_g (along with the associated uncertainties) on the remaining 731 candidate polymers. As shown in Fig. 1, we then use one of the three different strategies (explained in the following sections) to iteratively choose the next-best polymer to synthesize and test (from the 731 remaining polymers) so as to achieve our goal of designing a certain number of high- T_g polymers in the shortest possible time. Since this is a demonstrative problem, we do not actually synthesize and test the polymer but instead the T_g of this sixth

recommended polymer is “revealed” to the learning algorithm and added to the training set. A new ML model is trained on these six polymers and predictions are made for the remaining 730 polymers. This loop is repeated continuously until we have “discovered” 10 polymers possessing T_g greater than 450 K.

Although we only perform a “virtual” experiment in the current work, such frameworks can be used in synergy with actual experiments^[12] or time-consuming *ab initio* calculations.^[24,25] As mentioned earlier, we utilize one of the three strategies to determine the next-best experiment. The three strategies utilize exploitation (searching close to the area of the current best estimate), exploration (searching in unexplored areas), and balanced exploitation/exploration. As explained in more detail in the “Workflow” section, the exploitation strategy uses only predicted T_g to provide suggestions for the next-best polymer to synthesize and test. On the other hand, the exploration strategy provides suggestions based on just the uncertainty of predicted T_g of the candidate polymers (not the prediction itself). Finally, the balanced exploration/exploitation strategy provides a recommendation based on the utilization of both the predictions and the associated uncertainties of the remaining candidate polymers.

We statistically evaluate the aforementioned three strategies within the context of finding 10 polymers possessing T_g higher than 450 K. We also evaluate how the three different strategies perform when changing the size of the initial dataset and also as a function of the relative difficulty of the objective to be achieved.

Dataset

Data for this work were obtained from publicly-available collections of experimental measurements: *Polymer Handbook*,^[21] *Prediction of Polymer Properties*,^[22] and an online repository of polymer properties.^[23] The polymer dataset was highly-

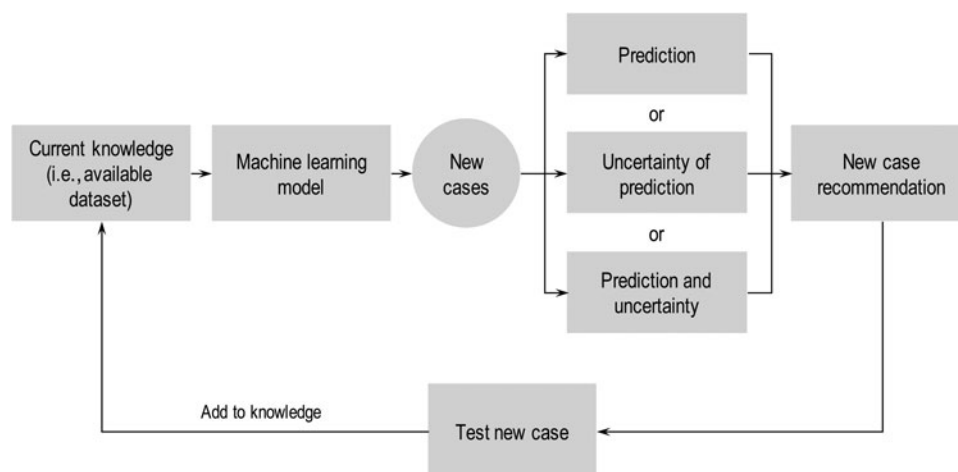


Figure 1. Overview of a typical active-learning framework. First, a model is trained based on the current knowledge of a system or environment. Using this model, predictions and associated uncertainties are obtained for new cases. Depending on the strategy that one wishes to employ, one may use the prediction, the uncertainty, or both the prediction and uncertainty to suggest the next-best case to be studied. Once the new case has been tested, the results thus obtained are used to update the current knowledge of the system and the iteration is repeated until a desired objective is achieved.

diverse and the constituent polymers were composed of nine atomic species: C, H, O, N, S, F, Cl, Br, and I. In order to visualize the diversity of 736 chemically unique polymers considered in this study, we performed principal component analysis (PCA) using 244 components of the hierarchical polymer fingerprint (described in the “Hierarchical polymer fingerprinting” section). Figure 2(a) shows distribution of polymers in a two-dimensional (2D) principal component space. Two leading components, PC1 and PC2 are assigned to x and y axes of the plot, respectively. The T_g of each polymer is used to color code the depicted points. As illustrated, the dataset not only includes common polymers such as polyethylene and polystyrene but also polymers with a large number of rings or those with very long side-chains. Also the T_g of the polymers in the dataset varied widely, ranging from 76 to 613 K with a mean of 326 K. A histogram of the distribution of T_g values in the dataset is shown in Fig. 2(b). The repeat unit of the polymers were represented using the simplified molecular-input line-entry system (SMILES).^[26]

Hierarchical polymer fingerprinting

In order to comprehensively capture the key features that may control the T_g , we utilized the hierarchical polymer fingerprinting scheme.^[27] The fingerprint building process consists of three hierarchical levels of descriptors. The first one is at the atomic scale wherein the occurrence of atomic triples (or a set of three contiguous atoms, e.g., C2–C3–C4, made up of a twofold coordinated oxygen, a threefold coordinated carbon, and a fourfold coordinated carbon) was calculated.^[28,29] For the polymers considered in this study, there are 123 such components. The next level deals with quantitative structure property relationship descriptors, such as van der Waals surface area,^[30] topological surface area,^[31,32] and fraction of rotatable

bonds,^[33] implemented in the RDKit cheminformatics library.^[34] Such descriptors, 99 in total, form the next set of components of our overall fingerprint. The third level and largest length scale descriptors captured morphological features such as the topological distance between rings, fraction of atoms that are part of side chains and length of largest side chain.^[27] We include a fixed set of 22 such morphological descriptors.

Workflow

As mentioned earlier, we set an arbitrary goal of finding 10 polymers possessing a T_g greater than 450 K. To achieve this objective, we would have to perform a series of virtual experiments, on one polymer at a time, to obtain the true experimental value of T_g for that polymer. At any given point of time, where the T_g of N polymers have been measured, how does one decide which polymer would be the best candidate for the $(N+1)$ th measurement? Instead of performing an experiment on a randomly selected polymer, it would be optimal if a suggestion was provided via an effective decision-making framework (leveraging experience gained from past measurements) to accomplish the goal of finding 10 high T_g polymers in the shortest possible time.

First, we start the design process by setting up an initial dataset with a random selection of five polymers from the entire dataset of 736 polymers. These five points are considered as “measured” polymers and will be used for training the initial surrogate model. Prior to developing the ML models we “fingerprinted” the polymers using a hierarchical polymer fingerprinting technique explained in the previous section. Among the fingerprint components, several morphological descriptors, such as the shortest topological distance between rings, fraction of atoms that are part of side-chains, and the length of the

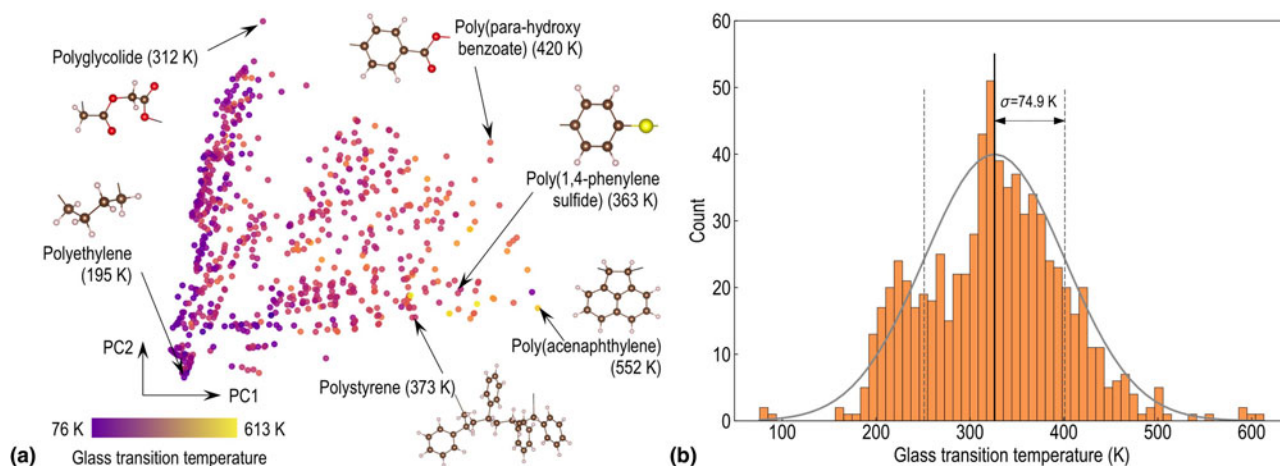


Figure 2. Graphical summary of the chemical space of polymers considered in this work. (a) 736 chemically unique polymers distributed in a 2D principal component space. Two leading components, PC1 and PC2 are produced by PCA, and assigned to x and y axes of the plot, respectively. The structure of a few representative polymers, with various number of rings and sizes of side chains are highlighted. (b) Distribution of the T_g values for all the polymers considered in this work.

largest side-chain, are included to properly capture the relevant features that could influence the T_g of a particular polymer.

Following the fingerprinting step, we created the surrogate model using GPR to learn the nonlinear relationship between the polymer fingerprints and their T_g values. A radial basis function kernel was utilized and fivefold cross validation was used to determine the hyperparameters of the new model at every iteration of the active learning workflow. Once the ML model has been built by training on the initial dataset, T_g values for the candidate polymers are predicted. We investigated the performance of three decision-making strategies (or acquisition functions); exploitation, exploration, and balanced exploitation/exploration, that employ different approaches (based on different criteria) to provide suggestions for the next-best experiment.

- **Strategy 1—Exploitation:** In this strategy the next-best polymer is the one that has the highest predicted T_g value. This prediction is obtained from an ML model trained on previous experiments. The exploitation approach favors polymers that are chemically similar to high- T_g polymers in the training set.
- **Strategy 2—Exploration:** Here, the next-best candidate is the one which displays the largest uncertainty (as predicted by the GPR model). Such an approach favors exploration of chemical space in order to acquire more information about the global landscape of T_g variation across all the polymers to be tested. As such, this strategy is not optimal for the targeted search of high- T_g polymers but we analyze the performance of this approach nonetheless.
- **Strategy 3—Balanced exploitation/exploration:** We utilize the maximum expected improvement (EI) acquisition function^[13,35] to utilize both the prediction and uncertainty to provide the suggestion for the next-best polymer candidate. For example, especially in the early stages of the dataset expansion, a particular polymer may not possess the highest predicted T_g but it may possess an unusually large uncertainty. Since the EI criterion takes into account both the prediction and the uncertainty, such a polymer may indeed be the most likely candidate for the next experiment. Within our GPR-based ML framework, the calculated EI metric is non-parametric and it automatically (and dynamically) provides a balance between exploration and exploitation approaches. The equations for the evaluation of the EI metric are shown in the Appendix.
- **Strategy 4—Random selection:** In this strategy, we randomly select a new candidate polymer, at every iteration, and add its T_g value to the list of “known” T_g values. The first three strategies are evaluated against this random approach.

In order to show the workflow in action, Fig. 3 demonstrates how the first three strategies utilize different metrics to choose the next-best polymer candidate for measurement at one particular iteration. Once 10 polymers possessing T_g greater than 450 K are obtained, the workflow is terminated and the number of iterations required to achieve this objective is noted. In order to remove the bias due to the initial (random) choice of five polymers, this entire workflow was repeated 50 times,

for each strategy, in order to obtain the average number of iterations required for the completion of the objective for each of the four strategies.

Since the performance of each strategy may vary depending on the size of the initial dataset or on the difficulty of the objective to be achieved, we also look into the relative performance of the four strategies when subject to the aforementioned variations. More specifically, we analyze the results when starting with different initial dataset sizes ranging from 5 to 60 and we also benchmark performance as a function of the threshold temperature (in the range of 300–450 K).

Results and discussion

As mentioned in the previous section, the above workflow is carried out to evaluate the performance of the three decision-making frameworks, i.e., exploitation, exploration, and balanced exploration/exploitation. In order to obtain statistically meaningful results, the workflow is repeated 50 times, each time starting with a different initial dataset (consisting of five randomly chosen polymers). Figure 4 demonstrates the number of experiments required (on average) to discover 1–10 polymers with a T_g of above 450 K. The error bars denote the standard deviation across the 50 different runs. For the purpose of comparison, the rates of success when using a random approach are also depicted in Fig. 4.

The average number of experiments required to discover 10 high- T_g polymers using the exploitation, exploration, balanced exploitation/exploration, and random approaches are 46, 98, 30, and 234, respectively. We now analyze the performance of each of the three different strategies individually.

- **Strategy 1—Exploitation:** The exploitation approach is the 2nd best-performing approach overall. From Fig. 4, we see that even though it has a very similar performance relative to the balanced exploitation/exploration approach, it possesses larger error bars. This implies that the exploitation strategy tends to get stuck in local minima, depending on the randomly chosen polymers in the initial dataset. As depicted in Fig. 5(a), we see that exploitation approach is the best-performing strategy when the initial dataset size is large. This is because larger initial dataset sizes lead to models that can predict T_g accurately over a wider range of chemical space. Also, as shown in Fig. 5(b), the exploitation approach also performs well when the threshold temperature for satisfying the selection criteria is lower. Lowering the threshold temperature reduces the difficulty of the search and exploitation approach is easily able to find high- T_g polymers in certain regions of chemical space.
- **Strategy 2—Exploration:** Of the three quantitative decision-making strategies, the exploration strategy is the least efficient one, even-though it still outperforms the random search approach. From Fig. 4, it is interesting to note that the exploration strategy has a sharper slope in the initial stages of the search but the rate of success slows down toward the later stages of the search. The exploration approach performs

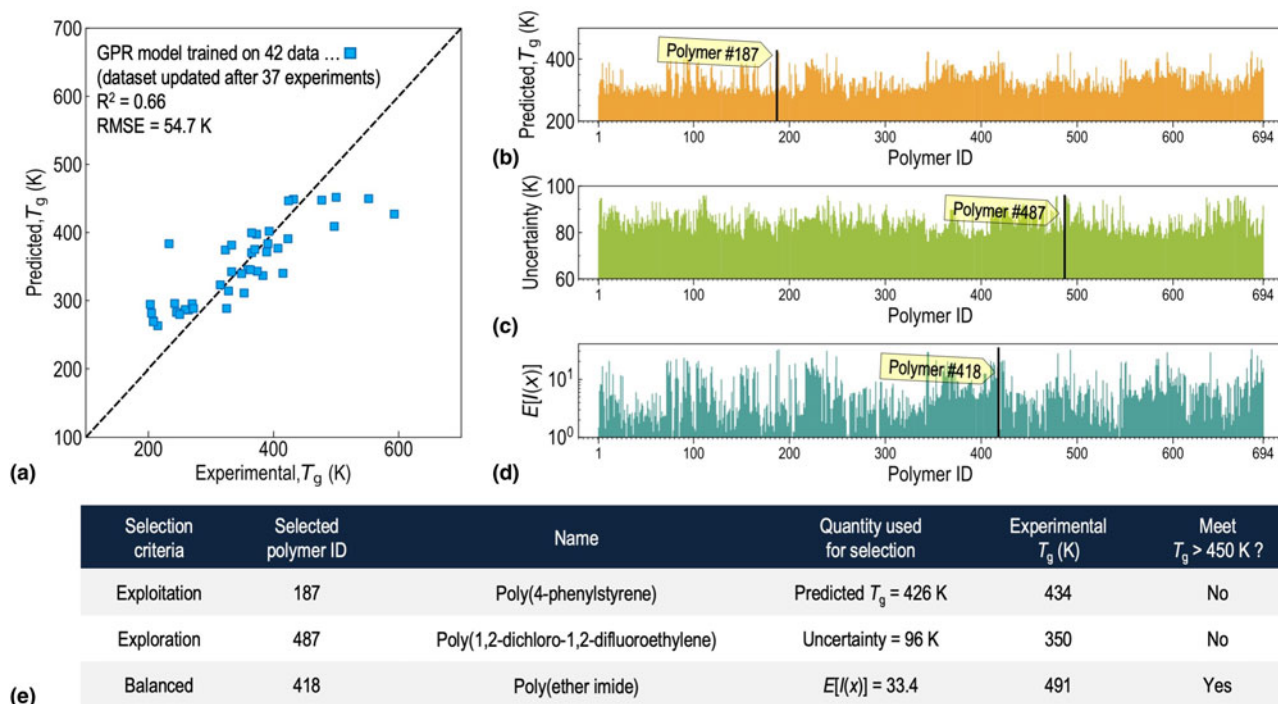


Figure 3. (a) Parity plot depicting ML predictions versus the actual T_g value for 42 polymers. Starting from an initial seed dataset of five polymers, 37 experiments were performed, resulting in the addition of 37 additional values to the dataset. The selection of the next-best candidate can be made using different criteria. (b) The predicted T_g values for the remaining 694 polymer candidates. Polymer #184 has the highest predicted T_g value and is thus selected when we use the exploitation approach. (c) Polymer #487 has the highest uncertainty associated with its prediction. When using the exploration framework, this particular polymer is the most suitable candidate since it indicates a point in chemical/fingerprint space that is least explored. (d) The EI metric is calculated for all remaining polymer candidates using both the polymer's T_g prediction and the corresponding uncertainty associated with the prediction. This balanced approach would lead to the selection of polymer #418 as the next-best candidate. In (e) we summarize the actual effect of making a selection based on the three different frameworks. In this particular case, the EI metric leads to the selection of the most suitable candidate, possessing a T_g which is above our required threshold of 450 K.

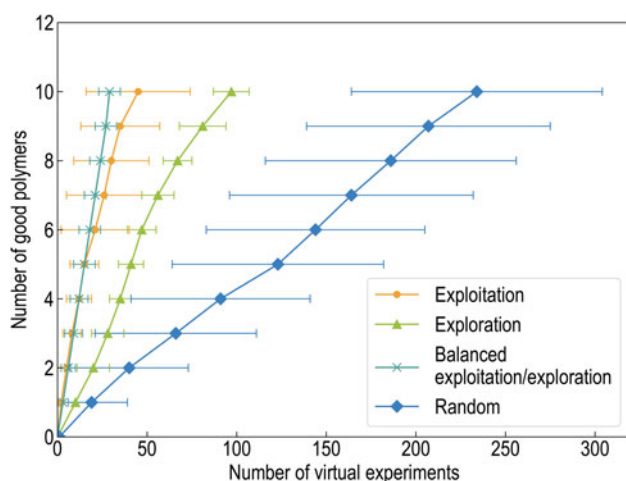


Figure 4. Number of experiments required (on average) to discover 1–10 polymers with T_g greater than 450 K when starting with an initial dataset size of five polymers. The average is calculated using 50 different runs and the standard deviation is denoted by the error bar.

well in the initial stages because it is able to explore a larger variety of chemical space as a result of using uncertainty to provide suggestions for the next-best test case. However, toward the later stages of the search, it is no longer required to search in unexplored regions since the ML model already has a good understanding of which regions of chemical space are likely to show high- T_g .

- **Strategy 3—Balanced exploitation/exploration:** The balanced exploitation/exploration approach is the best-performing and most robust strategy overall. As seen in Fig. 4, it not only performs better than the exploitation approach but it also shows small error bars and therefore it has a lower tendency to get stuck in local minima. It significantly outperforms the exploitation strategy when the initial size of the dataset is small. The EI metric dynamically provides the right balance between prediction and uncertainty of the prediction when suggesting the next-best polymer candidate to test. When the initial dataset size is small, uncertainties are larger across the test dataset and therefore this strategy will exhibit similar behavior as the uncertainty based exploration case. Toward the later stages of

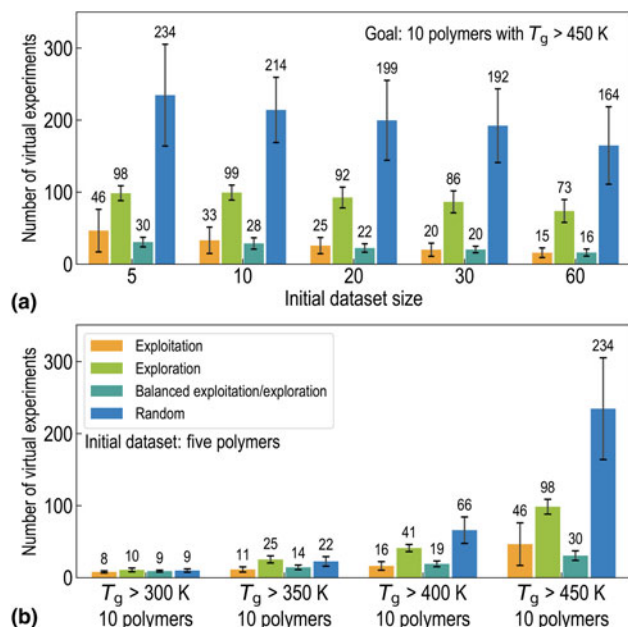


Figure 5. (a) The number of required experiments (on average) to find 10 polymers with T_g higher than 450 K starting from surrogate models trained on different initial dataset sizes. (b) Depicts the relative performance of the four different strategies on changing the threshold temperature for consideration as a high- T_g polymer. The black error bars denote standard deviations across 50 runs.

the search, the uncertainty (and the variation in uncertainty) is smaller across the remaining candidate polymers. Therefore, the balanced approach will perform in a manner akin to the exploitation strategy toward the end of search. For these reasons, as depicted in Fig. 5, the balanced approach performs exceptionally well when the initial dataset size is small and when the threshold temperature criteria is large.

- **Strategy 4—Random selection:** As expected, the random selection strategy is the most inefficient strategy for the discovery of high- T_g polymers relative to the other three decision-making strategies which are instead guided by quantitative metrics at every iteration. The random strategy only performs well when the threshold T_g temperature is low. This is simply due to the overall distribution of T_g values in the dataset which has a mean T_g value of 326 K as mentioned earlier.

Conclusion and outlook

To summarize, we have demonstrated an active-learning based approach that can be used to identify and discover new polymer materials possessing a high T_g . Within the context of Bayesian decision theoretic frameworks, we have evaluated the performance of three metrics that can be used to provide suggestions for the “next-best experiment.” We observed that the exploitation and balanced exploitation/exploration approaches (based on the EI criterion) showed the best performance in terms of rates of discovery of promising polymer candidates. The

balanced exploration/exploitation approach results in the most robust framework and performed especially well when the initial dataset size was small.

While ML techniques have been used widely in materials science over the past few years, such approaches are, more often than not, dependent on large amounts of data. In many cases the source of the data is a time-consuming experiment or quantum-mechanical calculation. While properties like T_g are widely available in the literature, other crucial properties such as dielectric breakdown^[8] or Li-ion diffusivity^[10] are hard to come by and difficult to measure. The integration of active-learning or BO frameworks within the materials discovery pipeline will provide quantitative guidance to systematically expand materials property datasets in an efficient and targeted fashion.

Acknowledgment

CK, AC and AJ were supported, respectively, by grants from the Office of Naval Research (Award Number N00014-16-1-2580), the Toyota Research Institute through the Accelerated Materials Design and Discovery program, and the National Science Foundation (Award Number 1743418).

References

1. A. Mannodi-Kanakithodi, T.D. Huan, and R. Ramprasad: Mining materials design rules from data: the example of polymer dielectrics. *Chem. Mater.* **29**, 9001–9010 (2017).
2. T.D. Huan, S. Boggs, G. Teyssedre, C. Laurent, M. Cakmak, S. Kumar, and R. Ramprasad: Advanced polymeric dielectrics for high energy density applications. *Prog. Mater. Sci.* **83**, 236–269 (2016).
3. A. Mannodi-Kanakithodi, G. Pilania, and R. Ramprasad: Critical assessment of regression-based machine learning methods for polymer dielectrics. *Comput. Mater. Sci.* **125**, 123–135 (2016).
4. T.D. Huan, A. Mannodi-Kanakithodi, C. Kim, V. Sharma, G. Pilania, and R. Ramprasad: A polymer dataset for accelerated property prediction and design. *Sci. Data* **3**, 160012 (2016).
5. A. Mannodi-Kanakithodi, G. Pilania, R. Ramprasad, T. Lookman, and J. E. Gubernatis: Multi-objective optimization techniques to design the pareto front of organic dielectric polymers. *Comput. Mater. Sci.* **125**, 92–99 (2016).
6. A. Mannodi-Kanakithodi, G. Pilania, T.D. Huan, T. Lookman, and R. Ramprasad: Machine learning strategy for accelerated design of polymer dielectrics. *Sci. Rep.* **6**, 20952 (2016).
7. A. Mannodi-Kanakithodi, A. Chandrasekaran, C. Kim, T.D. Huan, G. Pilania, V. Botu, and R. Ramprasad: Scoping the polymer genome: a roadmap for rational polymer dielectrics design and beyond. *Mater. Today* **21**, 785–796 (2018).
8. A. Mannodi-Kanakithodi, G.M. Treich, T.D. Huan, R. Ma, M. Tefferi, Y. Cao, G.A. Sotzing, and R. Ramprasad: Rational co-design of polymer dielectrics for energy storage. *Adv. Mater.* **28**, 6277–6291 (2016).
9. V. Sharma, C.C. Wang, R.G. Lorenzini, R. Ma, Q. Zhu, D.W. Sinkovits, G. Pilania, A.R. Oganov, S. Kumar, G.A. Sotzing, S.A. Boggs, and R. Ramprasad: Rational design of all organic polymer dielectrics. *Nat. Commun.* **5**, 4845 (2014).
10. D. Das, A. Chandrasekaran, S. Venkatram, and R. Ramprasad: Effect of crystallinity on Li adsorption in polyethylene oxide. *Chem. Mater.* **30**, 8804–8810 (2018).
11. S.P. Ong, O. Andreussi, Y. Wu, N. Marzari, and G. Ceder: Electrochemical windows of room-temperature ionic liquids from molecular dynamics and density functional theory calculations. *Chem. Mater.* **23**, 2979–2986 (2011).

12. M.K. Warmuth, J. Liao, G. Rätsch, M. Mathieson, S. Putta, and C. Lemmen: Active learning with support vector machines in the drug discovery process. *J. Chem. Inf. Comput. Sci.* **43**, 667–673 (2003). PMID: 12653536.
13. B. Shahriari, K. Swersky, Z. Wang, R.P. Adams, and N. De Freitas: Taking the human out of the loop: a review of Bayesian optimization. *Proc. IEEE* **104**, 148–175 (2016).
14. B. Rouet-Leduc, C. Hulbert, K. Barros, T. Lookman, and C.J. Humphreys: Automatized convergence of optoelectronic simulations using active machine learning. *Appl. Phys. Lett.* **111**, 043506 (2017).
15. R. Yuan, Z. Liu, P.V. Balachandran, D. Xue, Y. Zhou, X. Ding, J. Sun, D. Xue, and T. Lookman: Accelerated discovery of large electrostrains in BaTiO₃-based piezo-electrics using active learning. *Adv. Mater.* **30**, 1702884 (2018).
16. T. Mueller, A.G. Kusne, and R. Ramprasad: Machine learning in materials science: recent progress and emerging applications. In *Reviews in Computational Chemistry*, edited by A.L. Parrill and K.B. Lipkowitz (John Wiley & Sons, Inc., New York, **29**, 2016), pp. 186–273.
17. R. Ramprasad, R. Batra, G. Pilania, A. Mannodi-Kanakkithodi, and C. Kim: Machine learning in materials informatics: recent applications and prospects. *npj Comput. Mater.* **3**, 54 (2017).
18. D.J. Audus and J.J. de Pablo: Polymer informatics: opportunities and challenges. *ACS Macro Lett.* **6**, 1078–1082 (2017).
19. J.S. Peerless, N.J. Milliken, T.J. Oweida, M.D. Manning, and Y.G. Yingling: *Adv. Theory Simul.* **2**, 1800129 (2018).
20. S. Thrun: *Handbook of Brain Science and Neural Networks* (MIT Press, Cambridge, 1995), pp. 381–384.
21. J. Brandup, E.H. Immergut, and E.A. Grulke: *Polymer Handbook*, 4th ed. (John Wiley and Sons, New York, 1999).
22. J. Bicerano: *Prediction of Polymer Properties* (Marcel Dekker, Inc., New York, USA, 2002).
23. Polymer Properties Database. <http://polymerdatabase.com>, (accessed April 10, 2019).
24. B. Rouet-Leduc, K. Barros, T. Lookman, and C.J. Humphreys: Optimization of GaN LEDs and the reduction of efficiency droop using active machine learning. *Sci. Rep.* **6**, 24862 (2016).
25. L. Bassman, P. Rajak, R.K. Kalia, A. Nakano, F. Sha, J. Sun, D.J. Singh, M. Aykol, P. Huck, K. Persson, and P. Vashishta: Active learning for accelerated design of layered materials. *npj Comput. Mater.* **4**, 74 (2018).
26. D. Weininger: SMILES, a chemical language and information system. 1. Introduction to methodology and encoding rules. *J. Chem. Inf. Model.* **28**, 31–36 (1988).
27. C. Kim, A. Chandrasekaran, T.D. Huan, D. Das, and R. Ramprasad: Polymer genome: a data-powered polymer informatics platform for property predictions. *J. Phys. Chem. C* **122**, 17575–17585 (2018).
28. P. Pankajakshan, S. Sanyal, O.E. de Noord, I. Bhattacharya, A. Bhattacharyya, and U. Waghmare: Machine learning and statistical analysis for materials science: stability and transferability of fingerprint descriptors and chemical insights. *Chem. Mater.* **29**, 4190–4201 (2017).
29. T.D. Huan, A. Mannodi-Kanakkithodi, and R. Ramprasad: Accelerated materials property predictions and design using motif-based fingerprints. *Phys. Rev. B* **92**, 14106 (2015).
30. P. Labute: *J. Mol. Graph. Model.* **18**, 464–477 (2000).
31. P. Ertl, B. Rohde, and P. Selzer: Fast calculation of molecular polar surface area as a sum of fragment-based contributions and its application to the prediction of drug transport properties. *J. Med. Chem.* **43**, 3714–3717 (2000).
32. S. Prasanna and R. Doerksen: Topological polar surface area: a useful descriptor in 2D-QSAR. *Curr. Med. Chem.* **16**, 21–41 (2009).
33. K. Nguyen, L. Blum, R. van Deursen, and J-L. Reymond: Classification of organic molecules by molecular quantum numbers. *ChemMedChem* **4**, 1803–1805 (2009).
34. RDKit: Open Source Toolkit for Cheminformatics. <http://www.rdkit.org/> (accessed April 10, 2019).
35. A. Forrester and A.K.A. Söbester: *Engineering Design via Surrogate Modelling* (John Wiley and Sons, Chichester, West Sussex, 2008).

Appendix: Expected improvement

Balancing between exploitation and exploration requires the optimization of a constant A that controls the trade-off between exploitation and exploration in the statistical higher bound,

$$\text{HB}(x_i) = T_{\text{g,pred}}(x_i) - As(x_i) \quad (1)$$

where $T_{\text{g,pred}}(x_i)$ is the predicted T_{g} , and $s(x_i)$ is the uncertainty of the prediction given by the GPR model. x_i is the fingerprint vector of a given polymer i . $\text{HB}(x_i)$ becomes $T_{\text{g,pred}}(x_i)$ when $A = 0$ thus the selection of x_i will become pure exploitation. When $A = \infty$, $\text{HB}(x_i)$ is equivalent to $s(x_i)$ (pure exploration).

Instead of manually finding a good A value, $\text{HB}(x_i)$ can be determined by maximizing the EI which is given by,

$$E[I(x)] = P[I(x)][T_{\text{g,pred,max}} - T_{\text{g,pred}}(x)] + \frac{s(x)}{\sqrt{2}} \exp \left[- \left(\frac{T_{\text{g,pred,max}} - T_{\text{g,pred}}(x)}{2s(x)} \right)^2 \right] \quad (2)$$

$$P[I(x)] = \frac{1}{2} \left[1 + \operatorname{erf} \left(\frac{T_{\text{g,pred,max}} - T_{\text{g,pred}}(x)}{\sqrt{2}s(x)} \right) \right] \quad (3)$$

$P[I(x)]$ is the probability of improvement on selection of the next polymer with the fingerprint vector of x that will give an improvement on the best observed value so far, $T_{\text{g,pred,max}}$. Using this function, we no longer need to explicitly find the best A in Eq. (1). Detailed mathematical derivation can be found in Ref. 35.