



Cite this: DOI: 10.1039/d6dd00358c

Range-aware Bayesian optimization for discovering diverse designs within target property windows

Shengli Jiang,  Jason Wu, Charles M. Schroeder and Michael A. Webb *

In many materials and product design problems, desirable candidates exhibit properties that fall within an acceptable range rather than achieve a single optimum. Recovering multiple, distinct solutions that satisfy such specifications is also practically valuable, as some candidates may be preferred for reasons of cost, processability, stability, or other downstream considerations that are difficult to encode directly in an objective function. Here, we develop a range-aware Bayesian optimization (BO) framework, building on probability-of-feasibility ideas from constrained BO, in which the acquisition function directly scores the posterior probability that a candidate satisfies a target range. The framework naturally extends to parallel pursuit of multiple distinct specifications over a shared candidate space. Across benchmark tasks, range-aware acquisition generally recovers larger and more diverse sets of valid designs than standard BO baselines and recent goal-seeking methods. Its utility is further demonstrated in two practically motivated design case studies involving optimizing reaction conditions for polymer synthesis and sequence-defined oligomer discovery for prescribed optical absorption bands, supported by quantum chemical calculations. These results suggest that range-aware BO can provide a practical and sample-efficient foundation for specification-driven design, particularly when design flexibility and solution diversity are important considerations.

Received 9th June 2026
Accepted 22nd July 2026

DOI: 10.1039/d6dd00358c

rsc.li/digitaldiscovery

1 Introduction

Many materials and product design tasks are governed by specification satisfaction rather than single-property optimization or constraint. In these settings, a design is valid so long as its properties lie within prescribed ranges.^{1,2} We refer to these acceptable property intervals as “target ranges.” For example, polymers used in flexible electronics and aerospace composites are often designed around application-specific glass-transition windows;^{3,4} photovoltaic absorbers perform best when their band gaps are near values suited for spectral matching;^{5–7} and lubricants are viable within suitable viscosity ranges.^{8,9} This range-based perspective is also common in broader design practice, including Ashby-style materials selection, where property ranges guide candidate screening before objective-based ranking;¹⁰ pharmaceutical quality-by-design, where critical quality attributes are specified by acceptable ranges;¹¹ and product development, where needs are translated into target specifications with ideal and marginally acceptable values.¹² Because several candidates may satisfy the same prescribed range yet differ in synthesis, cost, stability, scalability, or processability, the utility of design campaigns can be enhanced should they return a portfolio of valid options rather than a single solution.^{13–15} Thus, the relevant search problem is often

not to optimize a single property but to recover, under limited evaluation budgets, distinct candidates whose properties fall within prescribed ranges. Recent machine-learning studies have improved prediction of pure-component, thermodynamic, and kinetic properties from molecular structure.^{16–18} Such structure–property models provide the surrogate layer needed for specification-driven screening, whereas the search problem addressed here begins once target ranges have been chosen.

Bayesian optimization (BO) is well suited to materials design because it enables sample-efficient exploration through probabilistic surrogate modeling.^{19–21} It has been applied successfully to reaction screening,^{22,23} alloy formulation,^{24,25} and polymer or protein design.^{26–29} Standard BO acquisition functions, including Expected Improvement (EI),²⁰ Probability of Improvement (PI),³⁰ and confidence-bound methods,³¹ are primarily designed to identify extrema rather than valid regions defined by specifications. Multi-objective BO methods^{32–34} extend this paradigm to Pareto-front discovery but still prioritize frontier points rather than candidates that lie within acceptable property ranges. Even when Pareto points happen to fall within a target range, interior valid designs that are dominated by frontier points are not preferentially recovered, although they may be equally desirable in practice. Consequently, conventional BO formulations are often misaligned with specification-driven design.

Recent work has extended BO toward goal-seeking and feasibility-aware tasks. Level-set estimation³⁵ and contour-

Department of Chemical and Biological Engineering, Princeton University, Princeton, NJ 08540, USA. E-mail: mawebb@princeton.edu

finding methods,^{36,37} including the Expected Feasibility Function (EFF),³⁸ prioritize threshold crossings or boundary structure rather than discovery of distinct candidates within prescribed ranges. Throughout, we refer to an acquisition sampling-based, if it draws posterior function samples and runs an auxiliary algorithm on them, and sampling-free, if it is computed directly from posterior summaries such as means, variances, and distribution functions. Bayesian Algorithm Execution (BAX) and its multi-target extension MultiBAX^{1,39} encode general target-seeking utilities through posterior sampling and algorithm execution. This can make the acquisition computationally demanding when many posterior samples or large candidate sets are required, and BAX may propose points outside the target range under high predictive uncertainty.² Target-constrained approaches such as t-EGO² focus sampling around a desired scalar target value but do not naturally address vector-valued specifications or multiple independent target ranges. A closely related approach in adaptive experimental design uses the probability of reaching a target range to plan parallel experiments, but only for a single output and a single target.⁴⁰ We extend this range-based view to vector-valued targets, to multiple target windows sharing one budget, and to the recovery of diverse valid designs. Constrained BO provides the closest conceptual connection through its probability-of-feasibility criteria.^{41,42} This family has also been extended to high-dimensional, large-scale, and robust constrained problems.^{43–46} The usual constrained-optimization objective is to locate an optimum that satisfies one or more feasibility constraints. Here, membership in a target window is itself the design objective, and the goal is to recover multiple valid candidates across one or more target windows. Related ideas appear in chance-constrained and robust optimization, which reason about the probability of satisfying constraints under uncertainty,^{46,47} and in novelty-search and quality-diversity methods, which reward diversity in behavior or design space.^{48–50} Both are complementary to our goal of recovering many candidates within specified output windows under a limited budget. Hierarchical scalarization methods such as Chimera⁵¹ let the user rank the objectives and assign each a tolerance, then fold them into a single objective function. Because the search is driven by this one aggregate score, such methods pursue the objectives in priority order rather than tracking whether each criterion lies within its target range. Together, these limitations motivate a BO framework in which range satisfaction is the primary search signal.

In this work, we develop a range-aware BO framework for specification-driven materials design, built around two range-aware, sampling-free acquisition functions, the Tolerance Ball (TB) and the Heaviside (HV). Unlike point-target methods, which preferentially sample near a specified value, or constrained BO approaches, which treat feasibility as a modifier of a separate objective, our framework uses the posterior probability of satisfying a target range as the acquisition objective itself. TB therefore uses the probability-of-feasibility idea but applies it to target-window discovery rather than constrained optimum seeking. It scores the posterior probability that a candidate lies within a multi-output tolerance window and

uses that probability directly as the acquisition. This pointwise score does not explicitly optimize diversity, but it can keep multiple feasible regions competitive when posterior probability remains high in each region. We compute the membership probability in closed form through a noncentral χ^2 approximation, so the acquisition is sampling-free and differentiable. Our contribution is the range-window formulation, its closed-form multi-output realization, and a systematic benchmark of diverse valid-design recovery. Because many design campaigns must pursue multiple specifications under a shared experimental or computational budget,^{26,29,52–56} including different product grades or operating windows within the same design space, our framework naturally extends to parallel target search across multiple target ranges. We compare the proposed acquisition functions against existing goal-seeking and target-driven methods based on their ability to recover diverse valid designs across target ranges of varying width and difficulty. Finally, beyond idealized benchmarks, we demonstrate the framework in two practical design settings. These include polymerization reaction condition design using kinetic Monte Carlo simulations and sequence-defined oligomer design for prescribed optical absorption bands supported by quantum chemical calculations.

2 Methods

2.1 Problem formulation

2.1.1 Single-target inverse design. We consider the problem of identifying designs $\mathbf{x} \in \mathcal{X} \subset \mathbb{R}^M$ whose K -dimensional property vector lies within a prescribed Euclidean tolerance of the target $\mathbf{y}^{\text{tgt}} \in \mathbb{R}^K$. Here, $\mathcal{X}$ denotes the design space, $\mathbb{R}^M$ is the M -dimensional real-valued space of possible designs, and $f: \mathcal{X} \rightarrow \mathbb{R}^K$ denotes the property map from a design to its predicted properties. A design $\mathbf{x}$ is considered valid if

$$\|f(\mathbf{x}) - \mathbf{y}^{\text{tgt}}\|_2^2 \leq \varepsilon^2, \quad (1)$$

where $\|\cdot\|_2$ denotes the Euclidean norm and $\varepsilon > 0$ is the tolerance radius. When target properties differ in scale or admissible tolerance radius, each output dimension is normalized before applying eqn (1). We use the Euclidean norm in the normalized output space because, after normalization, it treats deviations across output dimensions symmetrically and provides a simple measure of overall mismatch from the target. For a fixed target, we denote by $\mathcal{S} = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N\}$ the set of valid designs identified during the search, where $N = |\mathcal{S}|$ is the cardinality of $\mathcal{S}$ (*i.e.*, the number of members of the set). We denote by n_{eval} the total number of function evaluations allocated to that target.

2.1.2 Parallel search for multiple targets. We extend the single-target formulation to a collection of T targets, $\mathcal{Y}^{\text{tgt}} = \{\mathbf{y}_1^{\text{tgt}}, \dots, \mathbf{y}_T^{\text{tgt}}\}$. For each target $t \in \{1, \dots, T\}$, we then define the corresponding valid set by

$$\mathcal{S}_t = \left\{ \mathbf{x} \in \mathcal{X} : \|f(\mathbf{x}) - \mathbf{y}_t^{\text{tgt}}\|_2^2 \leq \varepsilon^2 \right\}, \quad (2)$$

and let $N_t = |\mathcal{S}_t|$ denote the number of valid designs identified for target t . Throughout and for convenience, we use a common

tolerance ε for all targets, although the formulation readily accommodates target-specific tolerances.

2.2 Performance metrics

We quantify search performance by the diversity of valid designs discovered per function evaluation. Diversity is used as an evaluation metric only, not as an explicit term in the acquisition functions.^{49,50}

2.2.1 Continuous design spaces. For continuous design spaces, diversity is measured using δ -uniqueness. A subset $S_\delta \subseteq S$ is δ -unique if

$$\|\mathbf{x}_i - \mathbf{x}_j\|_2^2 > \delta^2 \quad (3)$$

for all distinct $\mathbf{x}_i, \mathbf{x}_j \in S_\delta$.

Let $N_u(\delta)$ be the number of valid designs kept by a greedy δ -unique screen of S . The screen scans the valid designs in a fixed order and retains each one that lies farther than δ from all previously retained designs. We use this greedy count rather than the exact maximum δ -separated subset, which would require solving a combinatorial packing problem. We define the raw diversity score as the area under the uniqueness curve,

$$A_u = \int_0^{\delta_{\max}} N_u(\delta) d\delta, \quad (4)$$

where $\delta_{\max}$ is fixed at 0.1ϕ , with ϕ denoting the diameter of the design space. For the normalized design space in which each of the M design variables ranges from 0 to 1, $\phi = \sqrt{M}$. The normalized diversity score for continuous design spaces is then defined as

$$D_c = \frac{A_u}{n_{\text{eval}} \delta_{\max}}. \quad (5)$$

When performance is aggregated across multiple target tolerance ranges, we report the geometric mean of the corresponding D_c values.

2.2.2 Discrete design spaces. For discrete design spaces, the raw diversity is the number of distinct valid designs found, $N = |S|$. Let $N_{\text{valid}}^{\text{all}}$ denote the total number of valid designs in the full discrete candidate set. Because each function evaluation can identify at most one previously unseen valid design, the maximum attainable number of valid discoveries after n_{eval} evaluations is

$$N_{\max} = \min(n_{\text{eval}}, N_{\text{valid}}^{\text{all}}). \quad (6)$$

If $N_{\text{valid}}^{\text{all}}$ is unknown, we set $N_{\max} = n_{\text{eval}}$. We then define the normalized diversity score for discrete design spaces as

$$D_d = \frac{N}{N_{\max}}. \quad (7)$$

2.3 Gaussian process surrogate model

We model the mapping from candidates to properties using Gaussian process (GP) regression.⁵⁷ For a K -dimensional property vector, we fit one GP independently to each output dimension:

$$f_k(\mathbf{x}) \sim \mathcal{GP}(m_k(\mathbf{x}), \kappa_k(\mathbf{x}, \mathbf{x}')), k = 1, \dots, K, \quad (8)$$

using a zero prior mean, $m_k(\mathbf{x}) \equiv 0$. The zero prior mean is applied in the standardized output coordinate used by the benchmark. Each GP uses a Matérn 5/2 kernel with signal variance σ_k^2 and length scale l_k :

$$\begin{aligned} \kappa_k(\mathbf{x}, \mathbf{x}') \\ = \sigma_k^2 \left(1 + \frac{\sqrt{5}\|\mathbf{x} - \mathbf{x}'\|_2}{l_k} + \frac{5\|\mathbf{x} - \mathbf{x}'\|_2^2}{3l_k^2} \right) \exp\left(-\frac{\sqrt{5}\|\mathbf{x} - \mathbf{x}'\|_2}{l_k}\right). \end{aligned} \quad (9)$$

The kernel hyperparameters are re-estimated after each BO iteration by maximizing the marginal likelihood. The Matérn 5/2 kernel is a standard and robust default for Bayesian optimization;²¹ its twice-differentiable sample paths impose weaker smoothness than the squared-exponential (RBF) kernel⁵⁷ and thereby better accommodate the moderately rough structure-property relationships common in molecular data.⁵⁸

Given observations $\{(\mathbf{x}_i, y_{ik})\}_{i=1}^n$ for output k , the posterior predictive mean and variance at a new point $\mathbf{x}^*$ are

$$\begin{aligned} \mu_k(\mathbf{x}^*) &= \mathbf{k}_{k,*}^\top (\mathbf{K}_k + \sigma_{n,k}^2 \mathbf{I})^{-1} \mathbf{y}_k, \\ \sigma_k^2(\mathbf{x}^*) &= \kappa_k(\mathbf{x}^*, \mathbf{x}^*) - \mathbf{k}_{k,*}^\top (\mathbf{K}_k + \sigma_{n,k}^2 \mathbf{I})^{-1} \mathbf{k}_{k,*}, \end{aligned} \quad (10)$$

where $\mathbf{y}_k = [y_{1k}, \dots, y_{nk}]^\top$, $\mathbf{k}_{k,*} \in \mathbb{R}^n$ has entries

$$(\mathbf{k}_{k,*})_i = \kappa_k(\mathbf{x}^*, \mathbf{x}_i), \quad (11)$$

and $\mathbf{K}_k \in \mathbb{R}^{n \times n}$ has entries

$$(\mathbf{K}_k)_{ij} = \kappa_k(\mathbf{x}_i, \mathbf{x}_j). \quad (12)$$

Every oracle in this study returns a deterministic label. We keep the observation-noise term $\sigma_{n,k}^2$ in each GP fit as numerical regularization; in a physical experiment it would instead represent measurement noise. The TB probability defined below is therefore the posterior probability of target-range membership under the surrogate, not a guarantee of robustness to unmodeled variation.

2.4 Bayesian optimization procedure

Bayesian optimization is used to sequentially select candidates for evaluation. For a target $\mathbf{y}^{\text{tgt}}$, we define the squared discrepancy

$$d(\mathbf{x}) = \|f(\mathbf{x}) - \mathbf{y}^{\text{tgt}}\|_2^2 \quad (13)$$

and the posterior-mean discrepancy

$$\Delta^2(\mathbf{x}) = \|\boldsymbol{\mu}(\mathbf{x}) - \mathbf{y}^{\text{tgt}}\|_2^2. \quad (14)$$

Here $\boldsymbol{\mu}(\mathbf{x}) = [\mu_1(\mathbf{x}), \dots, \mu_K(\mathbf{x})]^\top$ stacks the per-output posterior means $\mu_k(\mathbf{x})$, so both discrepancies measure mismatch from the target $\mathbf{y}^{\text{tgt}}$ in the K -dimensional output space, with f the property map defined above.

For a single target, the BO loop proceeds as follows. First, the surrogate model is initialized with n_{init} candidates generated by Latin hypercube sampling (LHS).⁵⁹ The GP surrogate is then fit to the observed data, the acquisition function is maximized over $\mathcal{X}$, the selected candidate is evaluated, and the surrogate is updated. This process is repeated until the evaluation budget is exhausted. Unless otherwise noted, we use a total budget of 50 evaluations per run, and duplicate evaluations are not allowed.

For parallel target search, all targets are modeled by a single GP surrogate and a shared observation set. Search proceeds as a synchronous batch procedure. At the beginning of each outer iteration, the GP is fit once to all observed data. Each target then proposes one candidate by maximizing a target-specific acquisition function under the shared posterior. The resulting batch of proposed candidates is evaluated, and the GP is updated only after all evaluations in the batch are complete. The posterior therefore remains fixed while the batch is constructed. If two targets propose the same unevaluated candidate, that candidate is evaluated only once and assigned to the earlier target under a fixed ordering. The later target then selects the highest-ranked unevaluated alternative under the same posterior. This duplicate-removal step prevents redundant evaluation while preserving a batch size of T whenever sufficiently many unevaluated candidates remain. This procedure is target-parallel, proposing one candidate for each active target. For a single target with $q > 1$, TB can be used as the scalar score inside standard batch BO selection rules, such as posterior-fantasy selection or GP-BUCB.^{21,60}

2.5 Acquisition functions

We compare six acquisition strategies, Expected Improvement (EI), Lower Confidence Bound (LCB), the Tolerance Ball (TB) and Heaviside (HV) acquisitions, a posterior-mean BAX baseline (BAX), and random sampling (RS). We also report, in the SI, LCB across a range of exploration weights β , a constrained expected-improvement baseline that weights improvement toward the target by the probability of feasibility,^{41,42} a quality-diversity acquisition that also rewards spread in design space,^{48,49} and two sampling-based BAX variants, InfoBAX³⁹ and PS-BAX.⁶¹

2.5.1 Expected improvement. Expected Improvement (EI)²⁰ favors selecting candidates that are expected to improve upon the smallest discrepancy observed so far, with consideration also of the magnitude of that improvement. To obtain a closed-form expression for vector-valued targets, we adopt the isotropic predictive-variance approximation introduced by Uhrenholt *et al.*,⁶² treating the predictive distribution of $\mathbf{f}(\mathbf{x})$ given the current n observations $\mathcal{D}_n = \{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^n$ as:

$$\mathbf{f}(\mathbf{x})|\mathcal{D}_n \approx \mathcal{N}(\boldsymbol{\mu}(\mathbf{x}), \eta^2(\mathbf{x})\mathbf{I}_K), \quad (15)$$

where $\mathbf{I}_K$ is the $K \times K$ identity matrix and

$$\eta^2(\mathbf{x}) = \frac{1}{K} \sum_{k=1}^K s_k^2(\mathbf{x}) \quad (16)$$

is the isotropic predictive variance that averages the per-output posterior variances $s_k^2(\mathbf{x})$ of the independent GPs. Under this approximation,

$$\frac{d(\mathbf{x})}{\eta^2(\mathbf{x})} \sim \chi_K^2(\lambda(\mathbf{x})), \quad (17)$$

with noncentrality parameter

$$\lambda(\mathbf{x}) = \frac{\Delta^2(\mathbf{x})}{\eta^2(\mathbf{x})}, \quad (18)$$

where $d(\mathbf{x})$ is the squared discrepancy and $\Delta^2(\mathbf{x})$ is the posterior-mean squared discrepancy defined above. We then let

$$d_{\min} = \min_{i=1,\dots,n} d(\mathbf{x}_i) \quad (19)$$

denote the smallest observed discrepancy, and define

$$r = \frac{d_{\min}}{\eta^2(\mathbf{x})}. \quad (20)$$

The EI acquisition is then

$$A_{\text{EI}}(\mathbf{x}) = d_{\min} F_{K,\lambda}(r) - \eta^2(\mathbf{x}) [K F_{K+2,\lambda}(r) + \lambda(\mathbf{x}) F_{K+4,\lambda}(r)], \quad (21)$$

where $F_{K,\lambda}(\cdot)$ denotes the cumulative distribution function of the noncentral χ^2 distribution with K degrees of freedom and noncentrality parameter $\lambda(\mathbf{x})$.

2.5.2 Lower confidence bound. As a complementary baseline, we use the closed-form lower confidence bound (LCB) acquisition,⁶²

$$A_{\text{LCB}}(\mathbf{x}; \beta) = -(\alpha - \beta\rho)^{1/\ell} (K + \lambda(\mathbf{x}))\eta^2(\mathbf{x}), \quad (22)$$

where $\beta = 1$ controls the exploration–exploitation trade-off. The auxiliary quantities are

$$\begin{aligned} \ell &= 1 - \frac{r_1 r_3}{3r_2^2}, \\ r_s &= 2^{s-1}(s-1)!(K + s\lambda(\mathbf{x})), s = 1, 2, 3, \\ \alpha &= 1 + \ell(\ell-1) \left[\frac{r_2}{2r_1^2} - (2-\ell)(1-3\ell) \frac{r_2^2}{8r_1^4} \right], \\ \rho &= \frac{\ell r_2^2}{r_1} \left[1 - \frac{(1-\ell)(1-3\ell)}{4r_1^2} r_2 \right]. \end{aligned} \quad (23)$$

We use $\beta = 1$ for the main LCB results. SI Section S7 also reports fixed $\beta \in \{0.1, 0.3, 1, 3, 10\}$ and an adaptive choice following the finite-domain GP-UCB weighting,³¹

$$\beta_t = 2 \log \left(\frac{|C|\pi^2 t^2}{6\delta_u} \right), \delta_u = 0.1, \quad (24)$$

where t counts the evaluations selected after the initial design, with batch evaluations counted individually as in GP-BUCB.⁶⁰ For finite-pool tasks, $|C|$ is the number of candidates in the pool. For continuous tasks, where the optimizer is not a fixed finite-domain search, we set $|C| = 1$ as a documented heuristic rather than as a GP-UCB regret-bound setting.

2.5.3 Tolerance-aware acquisitions. We consider two tolerance-aware, sampling-free acquisitions computed directly from the GP posterior. The first, Tolerance Ball (TB), is the posterior probability that a candidate satisfies the output-space target range. Using the same isotropic predictive-variance approximation as in the EI derivation, we obtain

$$A_{\text{TB}}(\mathbf{x}) = \mathbb{P}[d(\mathbf{x}) \leq \varepsilon^2] = F_{K, \lambda(\mathbf{x})}\left(\frac{\varepsilon^2}{\eta^2(\mathbf{x})}\right). \quad (25)$$

For $K = 1$, this reduces to the probability that the scalar posterior lies within a tolerance interval around the target; for $K > 1$, it is the corresponding probability that the posterior vector lies inside a Euclidean ball. This is analogous to probability-of-feasibility criteria widely used in constrained Bayesian optimization,^{41,42,63} with feasibility here defined as membership in the prescribed output-space tolerance region. The Euclidean ball is the natural region when the outputs share a single tolerance budget, so that a small excess in one property can be offset by staying well inside the target on another. When each property must instead remain within separate lower and upper limits, the same probability-of-feasibility construction gives a box-shaped acquisition, whose membership probability is a product of per-output Gaussian interval probabilities because the outputs are modeled by independent GPs. A weighted ellipsoid is another direct extension, where a diagonal weighting matrix corresponds to property-specific tolerances after rescaling the outputs, whereas a full weighting matrix encodes correlated tolerances and leads to a Gaussian quadratic-form probability that can be evaluated numerically.

The second, Heaviside (HV), prioritizes candidates whose posterior mean lies inside the tolerance range while retaining a smooth transition near the boundary. We define

$$A_{\text{HV}}(\mathbf{x}) = (1 - w(\mathbf{x})) + w(\mathbf{x})A_{\text{TB}}(\mathbf{x}), \quad (26)$$

where

$$w(\mathbf{x}) = \frac{1}{2} \left[1 + \tanh\left(\frac{\Delta^2(\mathbf{x}) - \varepsilon^2}{\tau}\right) \right] \quad (27)$$

and $\tau = 10^{-3}\varepsilon^2$ is a smoothing parameter that controls the width of the transition region around the tolerance boundary. Thus, $A_{\text{HV}}(\mathbf{x}) \approx 1$ for candidates whose posterior mean lies well inside the tolerance range ($\Delta^2(\mathbf{x}) \ll \varepsilon^2$), whereas $A_{\text{HV}}(\mathbf{x}) \approx A_{\text{TB}}(\mathbf{x})$ for candidates whose posterior mean lies well outside it ($\Delta^2(\mathbf{x}) \gg \varepsilon^2$). At the boundary, $\Delta^2(\mathbf{x}) = \varepsilon^2$, so $w(\mathbf{x}) = 1/2$ and $A_{\text{HV}}(\mathbf{x}) = \frac{1}{2}(1 + A_{\text{TB}}(\mathbf{x}))$.

Notably, TB and HV are pointwise feasibility acquisitions. They do not explicitly penalize similarity to previously discovered valid designs. Any diversity in the discovered set therefore arises implicitly, from posterior uncertainty, the geometry of the feasible regions, and the exclusion of duplicate evaluations.

2.5.4 Posterior-mean BAX baseline. As a baseline for tolerance-constrained design discovery, we include a feasible-set exploration strategy based on a posterior-mean variant of BAX.^{1,39} Given that the corresponding feasible set of valid candidates is

$$\mathcal{T} = \{\mathbf{x} \in \mathcal{X} : d(\mathbf{x}) \leq \varepsilon^2\}, \quad (28)$$

we approximate^{1,39} this set by thresholding the GP posterior mean, resulting in

$$\overline{\mathcal{T}} = \{\mathbf{x} \in \mathcal{X} : \Delta^2(\mathbf{x}) \leq \varepsilon^2\}. \quad (29)$$

For continuous design spaces, $\overline{\mathcal{T}}$ is approximated on a set of 1000 candidate points sampled uniformly from $\mathcal{X}$ at each BO iteration. For discrete design spaces, the feasible-set approximation is evaluated over the full candidate set. The resulting BAX-style acquisition is defined as

$$A_{\text{BAX}}(\mathbf{x}) = \begin{cases} \frac{1}{K} \sum_{k=1}^K s_k(\mathbf{x}), & \text{if } \mathbf{x} \in \overline{\mathcal{T}} \setminus \mathcal{X}_{\text{obs}}, \\ 0, & \text{otherwise} \end{cases} \quad (30)$$

where $\mathcal{X}_{\text{obs}}$ is the set of previously evaluated candidates. In the first logic branch, ' $\setminus$ ' indicates the set-difference operator, such that the expression is evaluated over the set of candidates in $\overline{\mathcal{T}}$ that have not already been evaluated. If $\overline{\mathcal{T}} \setminus \mathcal{X}_{\text{obs}} = \emptyset$ (i.e., there are no unevaluated candidates in the target set), the acquisition defaults to global uncertainty sampling:

$$A_{\text{US}}(\mathbf{x}) = \frac{1}{K} \sum_{k=1}^K s_k(\mathbf{x}). \quad (31)$$

We refer to this baseline as posterior-mean BAX because it estimates the feasible set by thresholding $\Delta^2(\mathbf{x})$ rather than by sampling posterior functions. We also add two fully sampling-based variants that run the target-set algorithm on functions sampled from the GP posterior.^{64,65} InfoBAX selects the point whose evaluation is expected to be most informative about the sampled target sets;³⁹ we cap the posterior-sample count at 1, 3, 5, 10, or 15 so that its cost stays within a small multiple of the closed-form acquisitions. PS-BAX instead draws one posterior sample per iteration and evaluates the most uncertain member of its target set.^{61,66} Both variants follow the algorithms of the original works,^{39,61} and SI Section S9 gives the implementation details.

2.5.5 Random sampling. Random sampling selects the next unevaluated candidate at random from the available candidate set with uniform probability. For continuous design spaces, the candidate set consists of 1000 points drawn with uniform probabilities from $\mathcal{X}$ at each iteration. For discrete design spaces, the candidate set consists of all remaining unevaluated candidates.

2.6 Datasets

We evaluated the BO methods on four classes of inverse-design tasks. In each, the goal is to find candidates whose properties fall within a prescribed tolerance of a target specification: (i) synthetic benchmark functions (Branin, Hartmann-3, Ackley-5, Layeb-6) on continuous domains, for which the target is a scalar function value; (ii) pool-based datasets of pre-characterized candidates from nanoparticle synthesis, intrinsically disordered proteins, polymer topology, and small-molecule libraries, for which targets are measured or computed physicochemical properties; (iii) kinetic Monte Carlo (KMC) styrene polymerization, for which the design variables are reaction conditions and the target is the shape of the polymer molecular weight distribution; and (iv) sequence-defined conjugated oligomers, for which the design variables are molecular building-block

choices and the target is the shape of the UV-vis absorption band. In the following, we describe these design tasks and the underlying functions in more detail.

2.6.1 Synthetic functions. We considered four analytic benchmark functions, Branin ($\mathbb{R}^2 \rightarrow \mathbb{R}$), Hartmann-3 ($\mathbb{R}^3 \rightarrow \mathbb{R}$), Ackley-5 ($\mathbb{R}^5 \rightarrow \mathbb{R}$), and Layeb-6 ($\mathbb{R}^6 \rightarrow \mathbb{R}$). The analytic expressions and domains of all four functions are given in the SI Section S1. Fig. 1 shows representative two-dimensional slices with the target levels overlaid. These benchmark functions span a range of input dimensionalities and landscape characteristics, including differing degrees of multimodality and roughness, and they induce inverse-design tasks with multiple valid solutions that may lie in distinct regions of the design space.⁶⁷

2.6.2 Pool-based datasets. For pool-based experiments, the design space is a fixed finite candidate set. At each BO iteration, one previously unevaluated candidate is selected from the pool and evaluated. The datasets used were nanoparticle synthesis (NS), intrinsically disordered proteins (IDP), TopoRg, BACE, ESOL, FreeSolv, Lipophilicity (Lipo), and QM9.

SI Table S1 summarizes the number of candidates, the design-space representation, and the target property or

properties for each dataset. The NS dataset comprises 1997 candidate synthesis conditions, each defined by four normalized design variables, $\text{Ti}(\text{TEOA})_2$ concentration, TEOAH_3 concentration, pH, and temperature.⁶⁸ The target outputs are nanoparticle radius and polydispersity index. The IDP dataset includes 2031 polypeptide sequences of length 20–50 residues. Each sequence is characterized by three normalized phase-behavior properties, radius of gyration ($\bar{R}_g$), second virial coefficient ($\bar{B}_2$), and expenditure density (ρ^*).⁶⁹ The TopoRg dataset consists of 1340 candidates represented by eight-dimensional latent vectors obtained from a variational autoencoder trained on graph representations of polymer architectures. The target property is the mean single-chain radius of gyration.⁷⁰ The remaining datasets (BACE, ESOL, FreeSolv, Lipophilicity, and QM9) are molecular libraries with experimentally measured or computationally derived physicochemical properties.^{71–75} The specific target property used for each dataset is listed in Table S1.

2.6.3 Kinetic Monte Carlo styrene polymerization. We considered a task in which the objective is to identify reaction conditions that reproduce a prescribed molecular weight

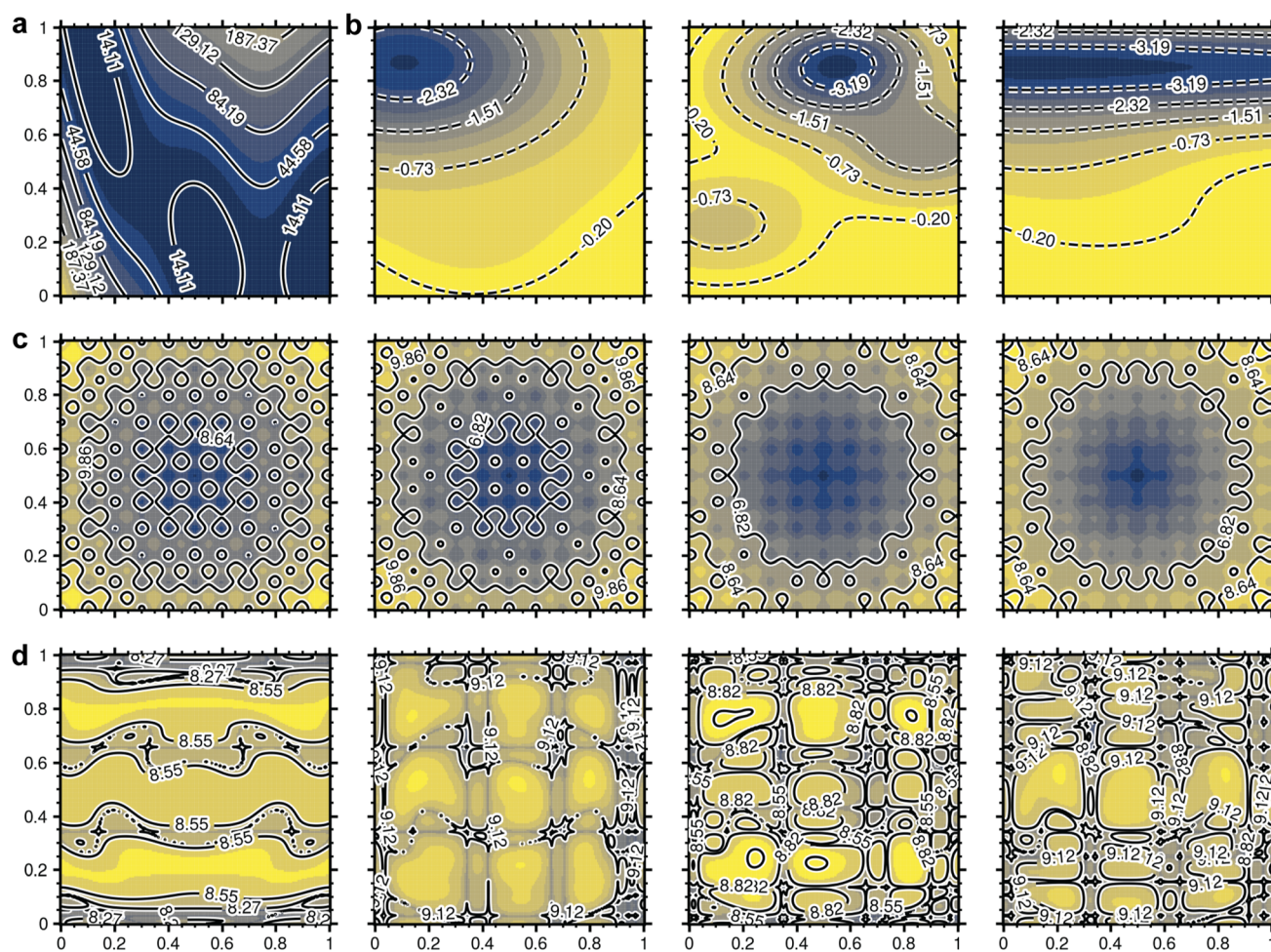


Fig. 1 Synthetic functions and targets. (a) Branin function ($\mathbb{R}^2 \rightarrow \mathbb{R}$). (b) Hartmann-3 function ($\mathbb{R}^3 \rightarrow \mathbb{R}$), shown as two-dimensional slices in which the first, second, and third coordinates (left to right) are fixed at 0. (c) Ackley-5 function ($\mathbb{R}^5 \rightarrow \mathbb{R}$), shown as two-dimensional slices over (x_1, x_2) with $x_4 = x_5 = 0$ and x_3 fixed at 0, 0.25, 0.5, and 0.75 (left to right). (d) Layeb-6 function ($\mathbb{R}^6 \rightarrow \mathbb{R}$), shown as two-dimensional slices along randomly selected coordinate pairs, with the remaining coordinates fixed at randomly selected values. Contours indicate the target levels in all panels.

distribution (MWD). To this end, we modeled styrene polymerization using a simplified two-stage radical polymerization process inspired by the sequential controlled/conventional strategy reported by Lenzi *et al.*⁷⁶ The first stage is a nitroxide-mediated polymerization (NMP) step, and the second stage is a conventional free-radical polymerization (FRP) step. This model is intended as a stylized batch, bulk, isothermal benchmark that captures the mechanistic contrast between controlled radical growth in stage one and conventional radical propagation in stage two, rather than as an exact reproduction of the heterogeneous reactor configurations used experimentally by Lenzi *et al.*^{76,77}

The complete reaction networks and kinetic parameters for both stages are listed in SI Tables S2–S4, using literature-based approximations for styrene NMP and FRP kinetics. The KMC scheme follows the stochastic simulation algorithm originally developed by Gillespie,^{78,79} with the polymerization formulation reported by Romero *et al.*⁸⁰ Each candidate is represented by a five-dimensional design vector

$$\mathbf{x} = (X_{\text{NMP}}, T_{\text{NMP}}, C_{\text{init,NMP}}, T_{\text{FRP}}, C_{\text{init,FRP}}), \quad (32)$$

where X_{NMP} denotes the prescribed monomer conversion at the end of the NMP stage, T_{NMP} and T_{FRP} are the reaction temperatures in the two stages, and $C_{\text{init,NMP}}$ and $C_{\text{init,FRP}}$ are the initiator concentrations in the NMP and FRP stages, respectively. Temperatures are reported in units of Kelvin (K), and initiator concentrations are reported in molar units (mol L^{-1}). The FRP stage was terminated at an overall conversion of 95%.

To generate the polymerization design space, we enumerated a full-factorial grid over these five variables. The prescribed NMP conversion X_{NMP} was varied from 0.4 to 0.8 using 8 uniformly spaced values. The NMP temperature T_{NMP} was varied from 353 to 403 K using 10 uniformly spaced values, and the FRP temperature T_{FRP} was varied from 313 to 363 K using 10 uniformly spaced values. The initial monomer concentration was 8.56 mol L^{-1} , and the KMC simulation was initialized with a control volume of $8.86 \times 10^{-16} \text{ L}$. The initiator concentrations $C_{\text{init,NMP}}$ and $C_{\text{init,FRP}}$ were determined from logarithmically spaced monomer-to-initiator ratios of 50 to 200 for NMP and 500 to 2000 for FRP, respectively, yielding 8 values for each stage. This corresponds to $C_{\text{init,NMP}}$ values from 4.28×10^{-2} to $1.712 \times 10^{-1} \text{ mol L}^{-1}$ and $C_{\text{init,FRP}}$ values from 4.28×10^{-3} to $1.712 \times 10^{-2} \text{ mol L}^{-1}$. The resulting full-factorial grid comprised 51 200 unique reaction conditions.

For a given candidate $\mathbf{x}$, the KMC simulation returns the final MWD in weight-fraction form. We represent this output as a 100-dimensional vector $\mathbf{y}$, where each component y_b is the polymer weight fraction within the b -th bin. The molecular weight of a chain with degree of polymerization n is $M = M_{\text{M}}n$, where $M_{\text{M}} = 104.15 \text{ g mol}^{-1}$ is the styrene monomer molar mass (SI Table S2). The bins are defined on a logarithmic molecular-weight axis spanning M_{M} to $10^5 M_{\text{M}}$ (approximately 10^2 to 10^7 g mol^{-1}), with 100 uniformly spaced intervals in log-10 space. Specifically, bin b corresponds to the interval

$$\log_{10}(M) \in [\xi_b, \xi_{b+1}), \quad (33)$$

where $\{\xi_b\}_{b=1}^{101}$ are uniformly spaced between $\log_{10}(M_{\text{M}})$ and $5 + \log_{10}(M_{\text{M}})$. The output vector is normalized such that

$$\sum_{b=1}^{100} y_b = 1. \quad (34)$$

This binned weight-fraction output serves as the property vector used for inverse design. We treat the 100-bin vector as a fixed numerical descriptor of the MWD rather than as a full probabilistic representation.

2.6.4 Sequence-defined conjugated oligomers. We considered a task defined over a library of sequence-defined conjugated oligomers. The library was constructed from 21 synthetically accessible oligomer cores and 100 commercially available pinacol boronic esters that serve as terminal caps. Combining each core with each cap yields 2100 candidate structures; after removing duplicate SMILES, which arise when distinct core-cap pairings are structurally equivalent, 1980 unique oligomers remain and constitute the optimization library (Fig. 6a and SI Fig. S15).

The objective is the shape of the simulated UV-vis absorption band. For each oligomer we computed a broadened absorption spectrum over 240–500 nm by time-dependent density functional theory (TD-DFT)^{81,82} and represented it by a compact, physically interpretable parametrization. Working in photon-energy space (eV), where the electronic bands are well described as Gaussian, we fit each spectrum with $K = 5$ Gaussian components,

$$S(E) = \sum_{k=1}^K A_k \exp \left[-\frac{(E - \mu_k)^2}{2\sigma^2} \right], \quad (35)$$

where A_k and μ_k are the amplitude and center energy of band k , and the width $\sigma = 0.12 \text{ eV}$ is held fixed for all bands and all molecules, matching the broadening used to construct the spectra. Spectra contain between one and five resolved bands, which we encode in a fixed-length representation by ordering the components by descending amplitude. When a spectrum has fewer than five bands, each unused component is assigned zero amplitude ($A_k = 0$) together with a single canonical center energy common to all molecules (3.35 eV). Because this canonical center is constant across the library, it has zero variance and is therefore uninformative to the surrogate, and it cancels exactly in the tolerance distance whenever both a target and a candidate lack the corresponding band. Moreover, $S(E)$ is independent of μ_k when $A_k = 0$, such that the center of an absent band has no effect on the absorption profile being matched; μ_k influences the objective only when band k is physically present ($A_k > 0$).

Each spectrum is peak-normalized such that its most intense band has unit amplitude ($A_1 \equiv 1$). Because A_1 is constant across the library, it is discarded. The resulting nine-dimensional optimization target ($\mu_1, A_2, \mu_2, A_3, \mu_3, A_4, \mu_4, A_5, \mu_5$) comprises the position of the dominant band together with the relative amplitudes and positions of the four secondary bands. This representation captures both where a molecule absorbs and the overall shape of its absorption envelope, allowing specifications

to be posed over an entire spectral band rather than over an isolated λ_{max} . Before optimization the target vector is scaled group-wise to $[0, 1]$, with all band amplitudes sharing one scaling and all band energies sharing another, such that the dimensionless amplitudes and the eV-scale energies contribute comparably to the tolerance distance rather than being normalized per dimension. The five resulting targets T_1 through T_5 are listed in SI Table S6 and displayed in Fig. S16.

2.6.5 Time-dependent density functional theory calculations. For each oligomer, long alkyl side chains were truncated to methyl groups to reduce computational cost,⁸³ explicit hydrogen atoms were added, and 100 three-dimensional conformers were generated in RDKit⁸⁴ using the ETKDGv3 embedding procedure,⁸⁵ each subsequently optimized with the MMFF94 force field;⁸⁶ the lowest-energy conformer was retained. To promote planarity of the conjugated backbone while preserving side-chain flexibility, dihedral-angle constraints (0° or 180° , chosen from the force-field-optimized conformer) were applied only to the single bonds linking distinct ring or fused-ring systems.⁸⁷ Constrained geometry optimizations were then performed in ORCA⁸⁸ using GFN2- $\times$ TB⁸⁹ with ALPB implicit solvation⁹⁰ in dichloromethane. Vertical excitation energies and oscillator strengths for the 15 lowest singlet excited states were computed for the optimized geometries by TD-DFT at the ω B97X-3c/CPCM(dichloromethane) level^{91,92} in ORCA. To obtain a continuous absorption profile rather than a discrete stick spectrum, only transitions above 240 nm were broadened (each by a Gaussian of width $\sigma = 0.12$ eV) and summed on a common grid spanning 240–500 nm, such that intense deep-UV bands below the cutoff did not contaminate the profile. The resulting broadened spectrum is the quantity subsequently parametrized by the five-Gaussian fit of eqn (35).

2.6.6 Input representation and preprocessing. All input variables were mapped to the unit hypercube $[0, 1]^M$ using min-max normalization performed separately for each dataset.

For the molecular datasets and the oligomer case study, each candidate was represented by the full set of molecular descriptors computed with Mordred from its chemical structure,⁹³ without manual feature selection. Descriptor columns with zero variance were removed, and principal component analysis (PCA)⁹⁴ was then performed separately for each dataset, retaining the number of components needed to explain 95% of the total variance; the surrogate therefore operates on these principal components rather than on individual descriptors. PCA is fit once, on the descriptor matrix of the fixed candidate set, before optimization begins; it is not refit as observations accumulate. Because this step uses descriptors only, it never sees the property measurements. The acquisition is evaluated only on candidates already in the pool, so every proposal is an existing candidate and no inverse-PCA decoding is needed.

For pool-based datasets, output properties were standardized to zero mean and unit variance separately for each property before target selection and tolerance definition. This standardization makes tolerance values comparable across datasets and across property dimensions. These benchmarks are retrospective, and the full candidate set is used to define the output

transform, the target medoids, and the tolerance scale. This is an oracle construction that provides a common coordinate system for comparing acquisition functions, not a deployable closed-loop protocol. In deployment, specifications would typically be defined in physical coordinates. The simplest approach is to fix the scaler in advance from calibration data, the initial design, or domain-specific property limits. If instead the scaler is refit as new observations are acquired, the mapping between physical and internal GP coordinates changes at each iteration; because the specification is fixed in physical space, the target bounds must be re-expressed in the updated coordinate system before evaluating the acquisition function. This does not affect the validity of the search, as the GP predictions and the physical-space specification remain consistent, but it does introduce additional bookkeeping relative to a fixed transform. SI Section S10 reports how target-window membership changes when the full-library scaler is replaced by one fit only on the initial design; the relative ranking of methods is preserved.

For analytic benchmark functions and case studies without a predefined output scale, no output normalization was applied.

2.6.7 Target and tolerance selection. For each dataset, we selected five representative targets that span the output space. For pool-based datasets, we applied k -medoids clustering⁹⁵ with $k = 5$ to the standardized property vectors of all candidates and used the resulting medoids as targets. For synthetic functions, we first sampled 20 000 inputs uniformly from the domain, evaluated the corresponding function values, and then applied k -means clustering⁹⁶ with $k = 5$ in output space. The resulting cluster centers were used as targets. We defined a dataset-specific base tolerance ε_0 as one half of the minimum pairwise Euclidean distances between the five targets. Performance was then evaluated at three tolerance levels, obtained by scaling this base tolerance by a tolerance ratio r ,

$$\varepsilon = r\varepsilon_0, r \in \{0.2, 0.4, 0.6\}, \quad (36)$$

where a larger tolerance ratio r yields a wider acceptance range and less stringent specification. In an experimental setting, the tolerance should also be set relative to the measurement and process-noise floor. A tolerance far below the noise scale makes observed range membership unstable, whereas one far above the relevant specification width can blur practically meaningful differences between candidates.

3 Results

3.1 Tolerance-ball acquisition achieves the best overall discovery performance

We first compare the acquisition functions on the benchmark suite, comprising the four analytical functions and the eight pool-based datasets, in terms of how effectively they recover multiple distinct valid designs under a fixed evaluation budget. This is the central criterion in our setting because specification-driven design often admits more than one acceptable solution, and a useful method should uncover a broad set of valid candidates rather than repeatedly refine a single promising

region. Accordingly, each task is scored by its normalized diversity score (D_c for continuous tasks, D_d for discrete tasks), and cross-task performance is summarized by ranks and the fraction of tasks where each method is best or worst at each tolerance ratio r .

Fig. 2 shows that TB exhibits the strongest discovery performance among the six main acquisition functions across the benchmark suite. It achieves the lowest average rank at all three tolerance ratios (Fig. 2a) and the highest fraction of best-performing tasks (Fig. 2b), ranging from 0.42 at $r = 0.4$ to 0.67 at $r = 0.6$. LCB and HV are close competitors, so we tested whether tuned or more directly feasibility-aware baselines change this interpretation. Sweeping fixed $\beta \in \{0.1, 0.3, 1, 3, 10\}$ and an adaptive GP-UCB schedule,³¹ no single fixed or adaptive LCB setting overturns the aggregate ranking, although choosing the best LCB variant separately for each setting would draw level with TB (SI Section S7). A quality–diversity acquisition, which rewards spread in design space under the same tolerance probability as TB, outperforms TB in 7 of the 36 settings and trails it in the other 29 (SI Section S8). The sampling-based BAX variants provide a stronger test than the posterior-mean BAX baseline. Across the 36 dataset–tolerance settings, TB ranks above the best capped InfoBAX in 31 settings and above PS-BAX

in 24, while PS-BAX ranks higher in 11 with one tie (SI Section S9). A constrained expected-improvement baseline performs on par with TB, attaining a slightly higher mean score while TB ranks first in more of the individual settings (19 of 36 against 16, with one tie). This parity is expected because the constrained-BO feasibility term is the same range-membership probability that TB scores directly (SI Section S11). We therefore read EI and RS as reference baselines, posterior-mean BAX as a mean-based set-estimation method, and LCB, constrained expected improvement, quality–diversity, InfoBAX, and PS-BAX as the stronger, more informative baselines. In contrast, RS performs worst overall and accounts for the majority of worst-performing outcomes (Fig. 2c), with EI and BAX also producing occasional last-place outcomes among the non-random methods.

The performance gap between TB and HV highlights the importance of retaining probabilistic resolution within and near the tolerance range. Both acquisitions are designed for targeting ranges of valid discoveries, but HV trails TB at every tolerance ratio. One likely reason is that HV assigns similar acquisition values to many candidates whose posterior means lie inside the range, even when those candidates have different posterior probabilities of satisfying the range. This behavior can

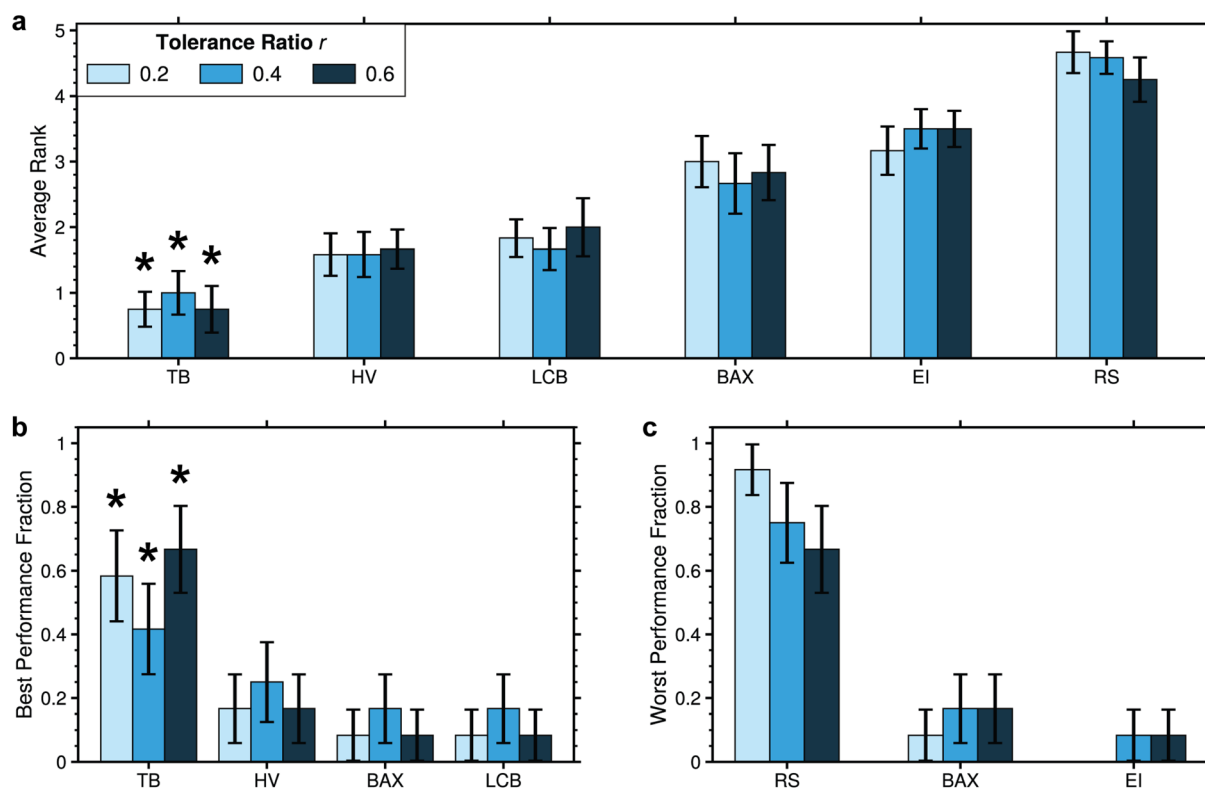


Fig. 2 Performance comparison of acquisition functions across tolerance ratios. (a) Average rank (lower is better) of each acquisition function across the benchmark tasks, including the four analytical functions and the eight pool-based datasets, at tolerance ratios $r = 0.2, 0.4, 0.6$. Within each task the six methods are ranked from 0 (best) to 5 (worst) by the normalized diversity score of that task, and these ranks are averaged over the twelve tasks. An asterisk marks the best-performing method at each tolerance ratio. (b) Fraction of tasks where each method achieved the best performance. (c) Fraction of tasks where each method achieved the worst performance. Methods that never attained the best (or worst) outcome are omitted from panel b (or c). The tolerance-ratio shading defined in panel a applies to all three panels. Bars show the mean across benchmark tasks; error bars show the standard error of the mean.

reduce its ability to prioritize the candidates with the highest probability of range satisfaction. This loss of resolution parallels the classical trade-off between Probability of Improvement (PI) and Expected Improvement (EI), in which PI saturates once a candidate is likely to improve and therefore disregards the magnitude of improvement that EI captures.⁹⁷ TB avoids this loss of resolution by directly scoring the posterior probability of range satisfaction, which provides a finer ranking of candidates according to their likelihood of meeting the specification.

The stronger performance of LCB relative to EI further indicates that extremum-seeking improvement is poorly aligned with range discovery. Although both methods use target discrepancy, EI evaluates improvement relative to the best discrepancy observed so far. This criterion is well suited to refining a single incumbent solution, but it can concentrate search near an already promising region after an acceptable candidate has been found. LCB combines posterior mean and uncertainty, allowing several plausible target-consistent regions to remain competitive during the search. This broader exploration behavior is more compatible with the goal of discovering multiple acceptable candidates, which likely explains why LCB consistently outperforms EI across tolerance ratios.

BAX underperforms relative to TB and LCB despite its set-oriented objective. This likely reflects its reliance on a posterior-mean approximation of the feasible region. When that approximation is incomplete, the acquisition may favor uncertainty reduction over candidates with high probability of satisfying the target range. Consequently, evaluations can be allocated to refining the feasible-set estimate rather than to recovering valid designs. This behavior is consistent with the weaker rank-based performance of BAX in Fig. 2. Drawing posterior function samples before building the target set, as InfoBAX and PS-BAX do, softens this concern but does not remove it on aggregate. Both sampling-based variants trail TB in the overall comparison, although PS-BAX ranks higher in 11 individual settings and is the strongest sampling-based alternative by that count (SI Section S9). Acting on a single sampled hypothesis each iteration can concentrate evaluations inside one plausible target region, which explains why PS-BAX is competitive in several settings even though it remains behind TB overall.

3.2 Tolerance-ball acquisition maintains balanced discovery across targets

Fig. 3 examines the per-target structure of discovery performance using the Branin benchmark and the QM9 molecular dataset as representative examples from the continuous and discrete settings, respectively; radar plots for the remaining benchmark tasks are provided in SI Section S4. The rank-based advantage of TB reflects its ability to sustain strong discovery across multiple targets rather than concentrating performance on only a subset. At an intermediate tolerance ratio ($r = 0.4$), TB achieves a high mean normalized diversity score D_c of 0.42 across the five targets in the Branin task (Fig. 3a), with comparable coverage across all five targets. HV and LCB also perform well in this low-dimensional setting, with mean scores of 0.41 and 0.42, essentially matching TB, whereas EI achieves a mean of 0.20 and

shows pronounced concentration on T_1 . BAX achieves an intermediate mean score of 0.36. A similar pattern appears in the QM9 inverse-design task (Fig. 3b), where TB again combines the highest mean score (0.15) with comparatively even performance across the full target set, while EI, BAX, and RS achieve lower scores and concentrate on a subset of targets. This balance is reflected in the across-target variance of D_c and D_d (SI Fig. S1). RS and EI show the largest across-target variation at moderate to high tolerance ratios, consistent with discovery concentrated on a subset of targets, whereas TB falls below both at $r = 0.4$ and $r = 0.6$. At $r = 0.2$, the methods are comparable. These results indicate that the overall performance of TB is accompanied by comparatively even discovery among targets.

The discovered-point plots further indicate that TB recovers candidates that are both valid and well distributed across the tolerance regions. In the Branin example (Fig. 3c), TB identifies candidates across multiple disjoint acceptable regions while keeping the discovered points closely confined to the prescribed tolerance ranges. HV also finds candidates in multiple regions, although its discoveries are less consistently confined to the admissible ranges. LCB often places points near the target contours, as expected from its level-set objective, and some of these fall just outside the tolerance regions. This reflects the difference between seeking a contour and scoring membership in the full range. BAX yields an even more diffuse distribution of discovered points, with a larger fraction of points outside the admissible regions.

The QM9 property-space visualization reveals the same qualitative pattern in a discrete molecular dataset. In Fig. 3d, TB produces a denser and more localized occupancy around the target locations, whereas LCB and HV generate more diffuse point clouds and BAX is the most dispersed. These visualizations support the interpretation that TB improves discovery by prioritizing candidates with high posterior probability of satisfying the tolerance criterion, rather than by broadly sampling near the target region. Additional radar plots in SI Section S4 show consistent behavior across the remaining benchmark tasks.

3.3 Off-target hit dynamics reveal target fidelity during parallel target search

In principle, algorithms for parallel target search should yield not only efficient discovery of valid candidates but also faithful allocation of those discoveries to their intended targets. In other words, the efficacy of an algorithm should not necessarily hinge on “accidental” discoveries where suitable candidates for one objective originate from the search for another. Fig. 4 tracks this property through the cross-target hit ratio. For each candidate proposed for a given target, we recorded whether the resulting property vector fell inside the tolerance range of another target. Because the target ranges considered here are disjoint, such an event indicates that the acquisition selected a valid design for the wrong target. The evolution of this ratio over BO iterations therefore provides a measure of target fidelity. Methods with low off-target hit ratios more consistently preserve the intended target, whereas methods with high ratios more frequently recover valid designs from other target ranges.

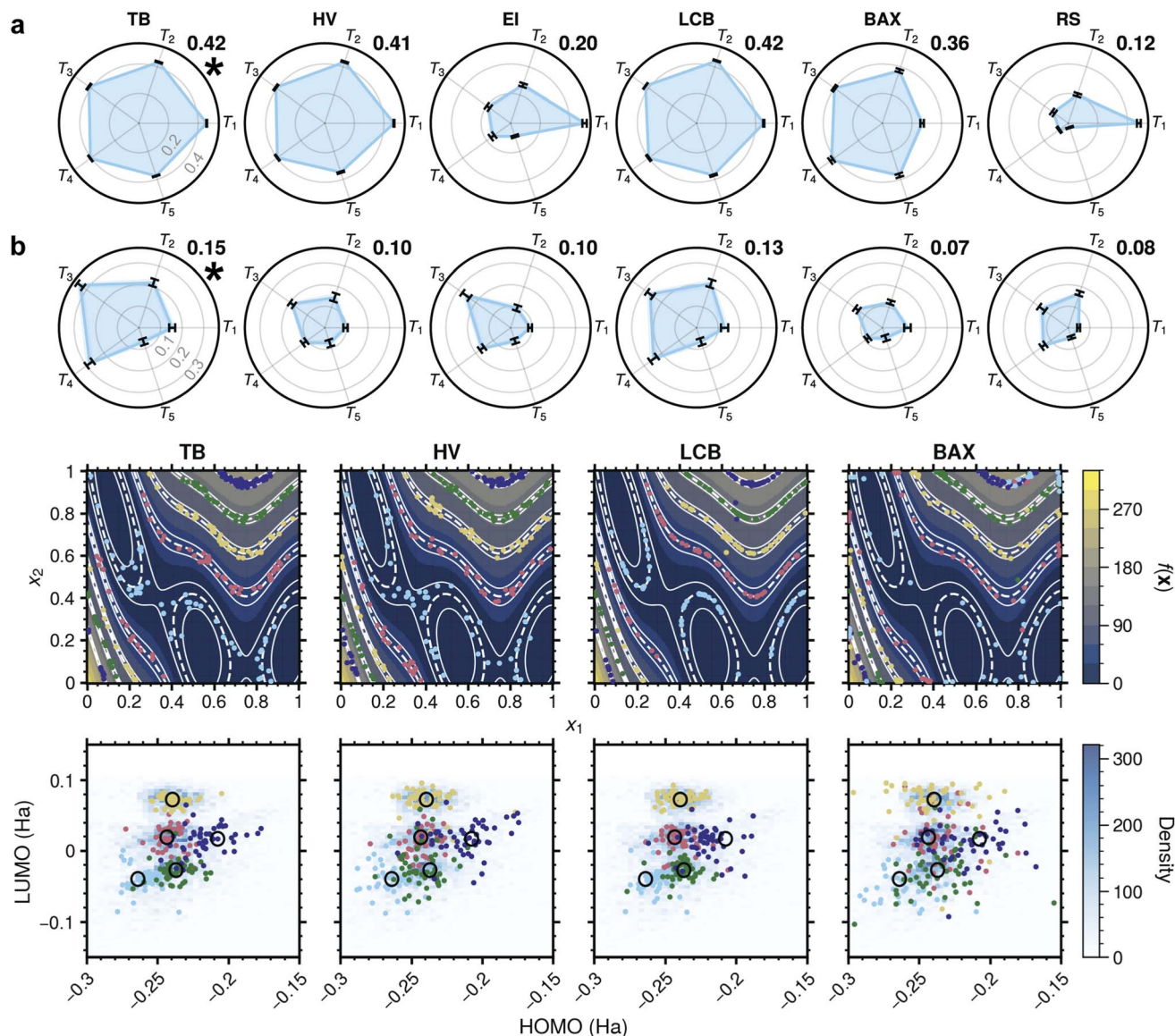


Fig. 3 Comparison of acquisition functions across the Branin and QM9 benchmarks. (a and b) Radar plots of the normalized diversity score (D_c for Branin, D_d for QM9) for five targets T_1 through T_5 in the (a) Branin and (b) QM9 tasks at $r = 0.4$, with the mean score reported on each plot. Each spoke corresponds to one target; larger radial values indicate higher normalized diversity scores, corresponding to a larger number of unique valid designs. Error bars indicate the standard error of the mean. An asterisk marks the best-performing method. (c) Comparison of solutions on the Branin function landscape. Background shading indicates the function value. White dashed contours denote the five target levels, and white solid contours denote the corresponding tolerance ranges. Colored markers indicate discovered points, with color denoting the associated target. (d) Comparison of solution on the QM9 property space defined by HOMO and LUMO. Background shading indicates the density of samples in the dataset. Open black circles mark the five target locations, and colored markers indicate discovered points with color denoting the associated target.

Target fidelity reflects both the acquisition function and the structure of the function landscape. In the Branin task (Fig. 4a), TB, HV, LCB, and BAX maintain near-zero off-target hit ratios across all tolerance ratios, whereas RS rises rapidly to a high plateau, which is expected since its search strategy is not actually targeted, and EI increases more gradually. The QM9 task (Fig. 4b) shows the same qualitative ordering between RS and the model-guided methods but the separation is smaller, and BAX and EI also drift upward at larger tolerance ratios. This behavior suggests that the design regions associated with different targets are less distinctly separated in the chosen

representation. The nanoparticle synthesis (NS) task (Fig. 4c) similarly preserves the advantage of TB, HV, and LCB, although BAX and EI here climb closer to RS. In the Lipophilicity (Lipo) task (Fig. 4d), all methods show more rapid growth in off-target hit ratio, indicating that the property landscape or candidate distribution makes target-specific discrimination intrinsically more difficult. Even in this more challenging setting, TB and LCB attain the lowest off-target hit ratios.

Higher off-target hit ratios are associated with weaker target-specific discovery in these experiments. RS produces the highest off-target hit ratios and also shows the weakest overall discovery

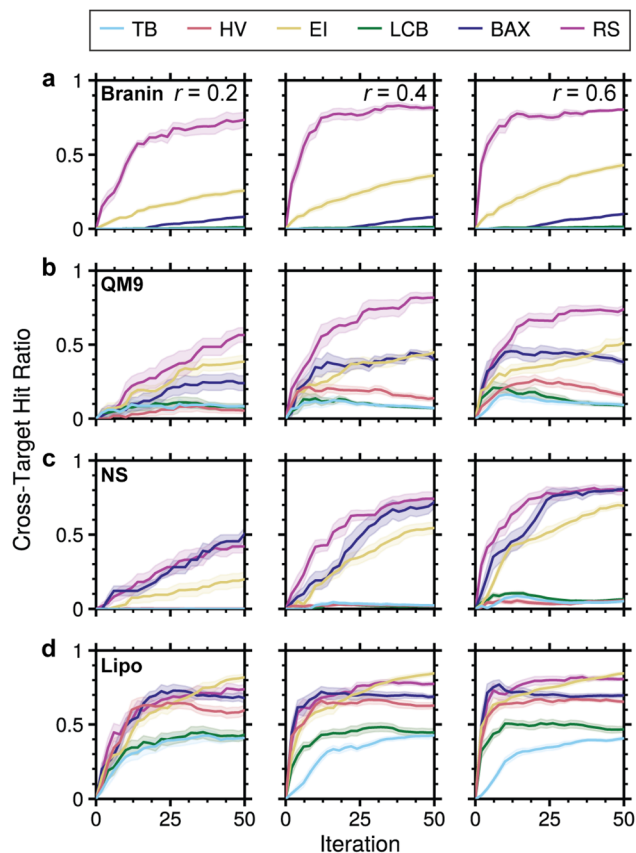


Fig. 4 Off-target hit dynamics during parallel target Bayesian optimization. Cross-target hit ratio as a function of BO iteration for the (a) Branin, (b) QM9, (c) Nanoparticle synthesis (NS), and (d) Lipophilicity (Lipo) tasks at tolerance ratios $r = 0.2, 0.4, 0.6$. Solid lines show the mean across repeated runs. Shaded regions show the standard error of the mean.

performance, whereas TB combines the lowest off-target hit ratios with the strongest rank-based discovery performance (Fig. 2). This relationship suggests that the advantage of TB is unlikely to arise from opportunistically recovering valid designs from any target range but instead from selecting candidates with high posterior probability of satisfying the intended target range. The off-target analysis therefore provides additional evidence that explicit range-satisfaction scoring can improve both discovery efficiency and target fidelity in parallel target search.

3.4 Case study: polymer molecular weight distribution design

Whereas the preceding benchmark tasks are drawn from existing datasets, we now turn to case studies designed to more closely reflect real-world considerations. The KMC polymerization task addresses inverse polymer design for a prescribed full molecular weight distribution (MWD). Rather than optimizing a scalar property such as number-average molecular weight or dispersity, the objective here is to identify reaction conditions that reproduce a target MWD. Designing the full MWD is well motivated, because its breadth, skew, and modality influence polymer properties beyond what summary statistics capture.^{98,99} This goal has been pursued through controlled-synthesis and flow

strategies^{100–103} and more recently through machine-learning approaches, including reinforcement learning and KMC-based active learning and multi-objective reverse engineering.^{104–108} However, recovering reaction conditions for a target MWD under an explicit, range-aware tolerance criterion remains comparatively less explored. Here, each specification is a 100-bin weight-fraction distribution, and a design counts as valid when its simulated distribution falls within the prescribed Euclidean tolerance of the target. Because the design space is substantially larger (51 200 candidates) and valid regions are sparser than in the pool-based benchmarks, we extended the evaluation budget to 200 evaluations for this case study. The five target distributions span a range of shapes, including unimodal profiles centered at different molecular weights and bimodal profiles with a dominant peak and a weaker high-molecular-weight shoulder (Fig. 5a). Fig. 5b also shows that the residuals between target and recovered distributions remain small in magnitude across the full molecular-weight range, indicating that the Euclidean tolerance criterion captures the principal features of the distribution in this fixed representation.

TB provides the strongest and most balanced recovery of valid polymerization conditions. In the radar plot of normalized diversity score (Fig. 5c), TB achieves the highest mean D_d of 0.93 across the five targets, followed by HV at 0.85, LCB at 0.69, and BAX at 0.64. EI and RS perform poorly, with mean scores of 0.10 and 0.04 respectively, reflecting that valid reaction conditions occupy only a small region of the five-dimensional design space and are unlikely to be recovered by extremum-seeking or unguided search. The target-resolved discovery trajectories further show that TB identifies valid designs more rapidly and consistently than the other methods. Fig. 5d reports the cumulative number of valid designs found for each target over 200 BO iterations. TB rises steeply during the early stage of the search and reaches the largest, or near-largest, number of valid designs for most targets by the end of the evaluation budget. LCB and HV also accumulate valid designs at competitive rates, whereas BAX trails the leading methods. EI and RS remain near zero throughout the search.

The valid designs discovered by TB are also diverse in reaction-condition space, even when they produce similar MWDs. For each target, a representative pair of valid designs that yield nearly identical distributions is reported in SI Fig. S14; the two designs differ across all five reaction variables (SI Table S5), confirming that distinct combinations of conditions can satisfy the same product-level specification. The UMAP¹⁰⁹ projection in Fig. 5e illustrates this, with the valid designs for every target occupying multiple separated regions of the embedding. This degeneracy is valuable for process design because it reveals multiple reaction pathways to the same product specification and provides flexibility for subsequent optimization of cost, stability, scalability, or experimental feasibility.

3.5 Case study: sequence-defined conjugated oligomers discovery

The sequence-defined oligomer task tests range-aware discovery in a finite molecular library, where discrete changes in core

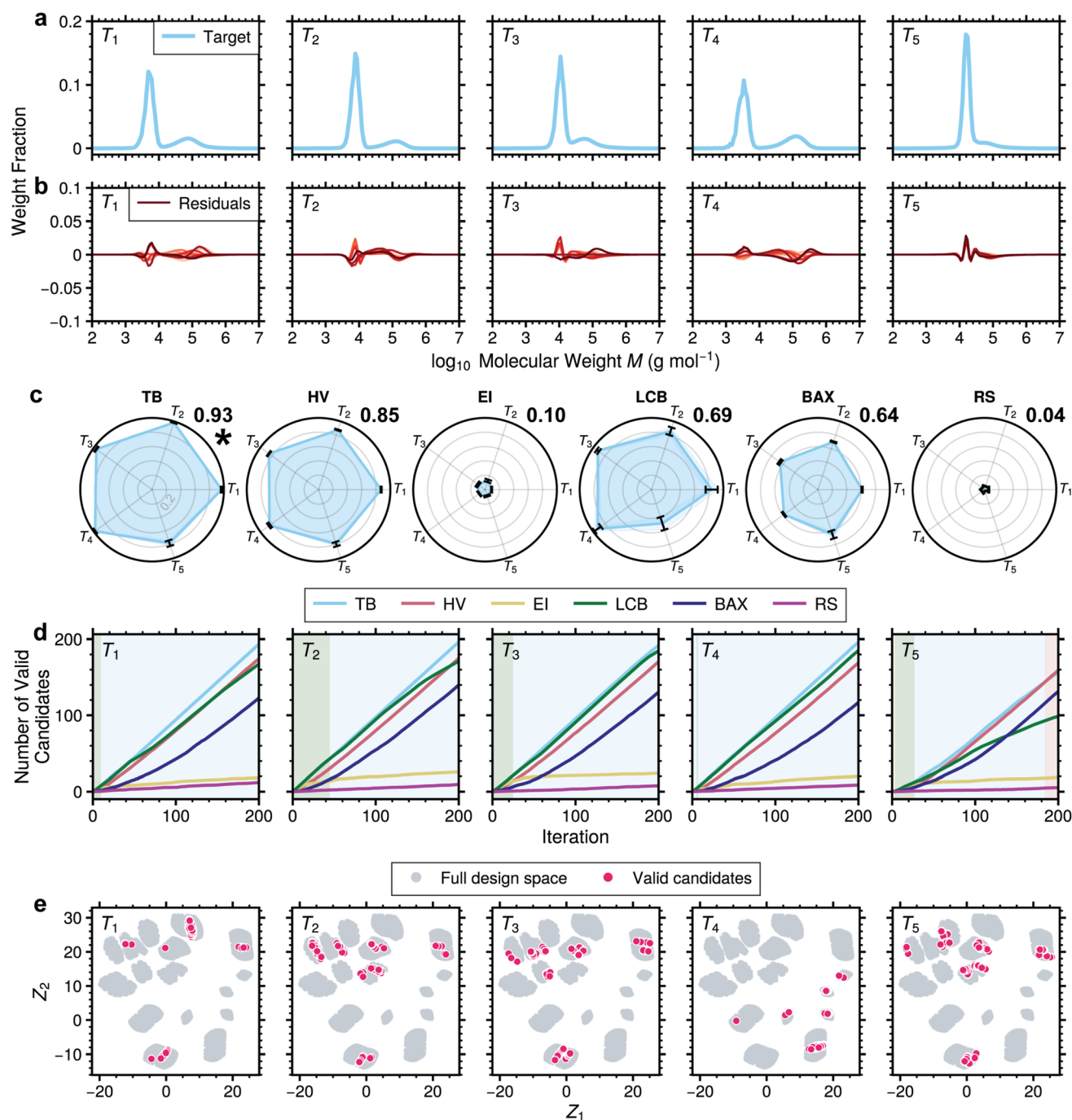


Fig. 5 Polymer molecular weight distribution design. (a) Five target MWDs and (b) the corresponding residuals between target and recovered distributions for valid designs discovered by TB, plotted as a function of $\log_{10}$ molecular weight. (c) Radar plot of normalized diversity score D_d across the five targets T_1 through T_5 for all acquisition functions, with the mean score reported on each plot and error bars indicating the standard error of the mean. An asterisk marks the best-performing method. (d) Cumulative number of valid designs identified as a function of BO iteration for each target T_1 through T_5 , with shaded background indicating the dominant acquisition function for each target. (e) UMAP projection of the five-dimensional reaction-condition space and the valid designs discovered for each target T_1 through T_5 , with axes Z_1 and Z_2 denoting the two UMAP components.

sequence or terminal-cap chemistry reshape the entire optical absorption band rather than merely shifting its peak. Accordingly, each target is specified not by a single wavelength but as a region of the spectral-shape space introduced above, that is, the nine-dimensional target vector ($\mu_1, A_2, \mu_2, A_3, \mu_3, A_4, \mu_4, A_5, \mu_5$) comprising the dominant-band energy together with the

relative amplitudes and positions of the four secondary bands; a representative spectrum and its five-Gaussian decomposition are shown in Fig. 6b. The residual between the simulated spectrum and this five-Gaussian fit (Fig. 6b, bottom) remains small and largely structureless, confirming that the compact parametrization reproduces the absorption band. The five

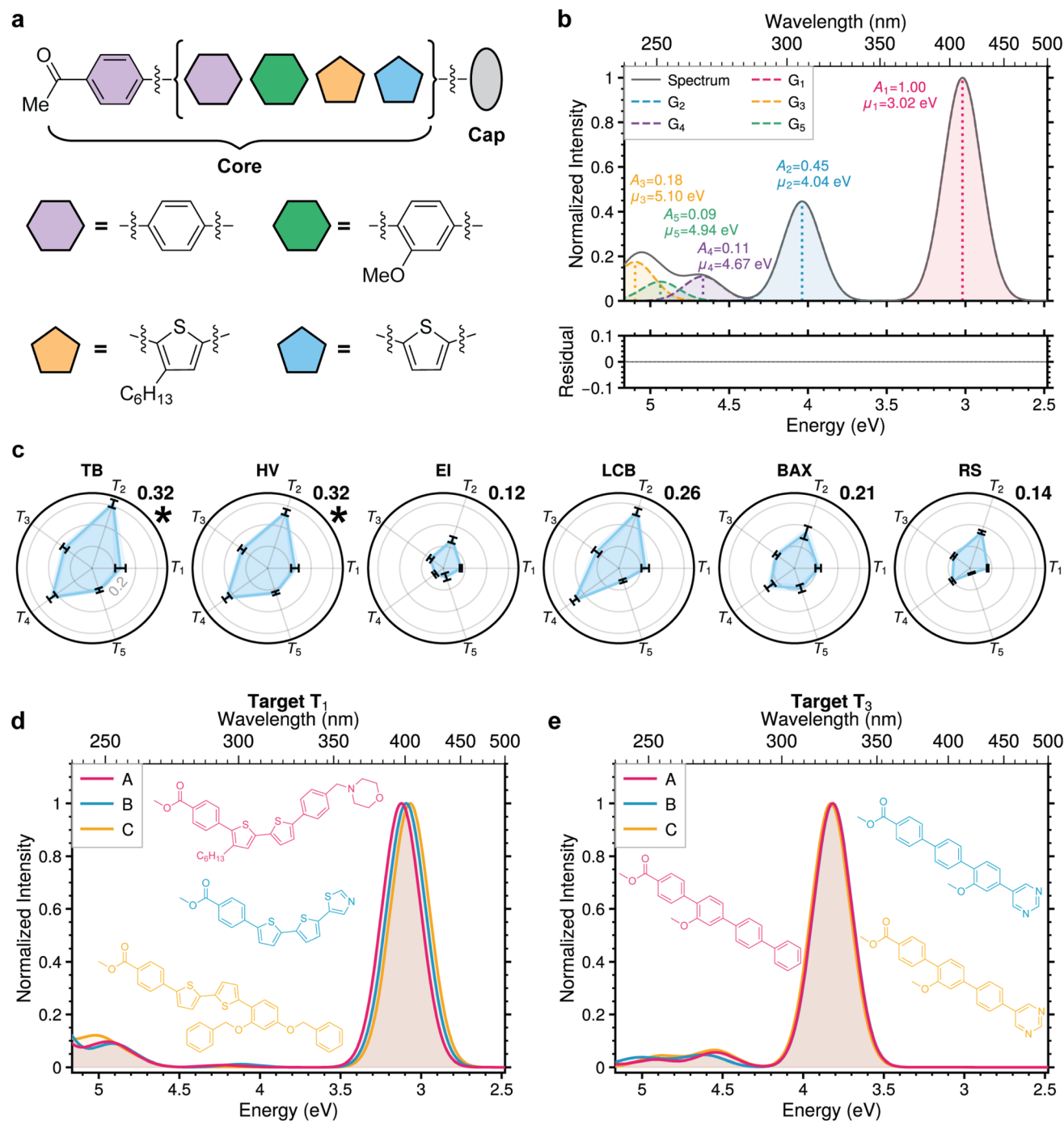


Fig. 6 Tolerance-aware discovery of sequence-defined conjugated oligomers. (a) Schematic of the oligomer architecture, comprising a methyl ester head group, a conjugated core assembled from four monomer types, and a terminal cap. (b) Representative simulated absorption spectrum decomposed in energy space into five Gaussian components G_1 – G_5 , each parametrized by an amplitude A_k and center μ_k (with fixed width); the fit residual is shown below and the top axis gives the corresponding wavelength. (c) Radar plot of normalized diversity score D_d across the five targets T_1 through T_5 for all acquisition functions, with the mean score reported on each plot and error bars indicating the standard error of the mean. An asterisk marks the best-performing method. (d and e) Three chemically distinct oligomers (A–C) discovered by TB for targets T_1 (d) and T_3 (e), whose simulated spectra are similar.

targets T_1 – T_5 were obtained by k -medoids clustering of the Gaussian spectral vectors and span the dominant-band variation observed in the library (see SI Table S6 and Fig. S16).

Performance differences among acquisition functions are less pronounced in this discrete molecular setting than in the

polymerization case study (Fig. 6c). TB attains the highest mean normalized diversity score ($D_d = 0.32$) and remains among the strongest methods across the five target bands, but the separation among methods is modest and random sampling is competitive for several targets. The target-resolved discovery

trajectories in SI Fig. S17 corroborate this view. From this data, our observation is that no single acquisition function dominates all bands across the full evaluation budget, although TB and HV sustain higher discovery rates.

The TB discoveries nevertheless reveal a chemically important form of degeneracy, where structurally diverse oligomers can lead to comparable absorption profiles. Fig. 6d and e present three TB-discovered oligomers for targets T_1 and T_3 , respectively. Within each panel the molecules differ in conjugated-core sequence and terminal-cap chemistry, yet their simulated spectra are similar, sharing an almost identical dominant-band position to within the prescribed tolerance. This many-to-one mapping from molecular structure to optical response is valuable from a design perspective because it provides multiple synthetic routes to a single absorption specification, leaving room for subsequent selection according to criteria outside the present objective, such as synthetic accessibility or stability.

4 Conclusion

Many materials and product design problems are guided by specification satisfaction. A candidate is valuable when its properties fall within prescribed tolerance ranges, rather than when a single objective reaches an optimum. We reformulated this challenge as a parallel target-range discovery problem and derived a closed-form Tolerance Ball (TB) acquisition that directly scores the posterior probability of range satisfaction through a noncentral χ^2 approximation, applying probability-of-feasibility ideas from constrained BO to multi-output target windows. Because the acquisition is sampling-free and analytically tractable, it offers a computationally efficient alternative to extremum-seeking strategies and sampling-based set-estimation methods.

Across synthetic functions, molecular libraries, and case studies, TB achieved the strongest aggregate recovery among the six main acquisition functions. LCB was a close second and could match or exceed TB on individual tasks, and constrained expected improvement was comparable in the supplementary comparison. This performance likely arises because TB explicitly scores range satisfaction for the intended target, yielding low cross-target hit ratios and high target fidelity. Improvement-based or uncertainty-driven methods, in contrast, can drift toward off-target ranges or concentrate on refining a single incumbent solution. In the large, sparse design space of the kinetic Monte Carlo polymerization task, TB recovered distinct reaction conditions that accurately reproduced prescribed molecular weight distributions (MWDs). This distributional target class has practical relevance for optimizing mechanical, rheological, and processing behavior. A similar structure-property degeneracy emerged in the sequence-defined oligomer library, where chemically distinct molecules produced nearly identical absorption bands, providing multiple molecular candidates for a single optical specification.

TB is not the top-ranked method on every individual target or dataset, but it performs strongly across the full range of settings, which makes it a reliable default for specification-

driven search when the shape of the feasible region is not known in advance. How much surrogate guidance helps also depends on the design space. When a discrete library already contains many valid molecules, the gap between TB and random sampling narrows, so model guidance adds the most value when valid regions are sparse or disconnected. Future work could incorporate explicit diversity incentives into the acquisition, relax the isotropic variance approximation through full multi-output Gaussian process models, model aleatoric or input uncertainty for robustness-aware specifications, replace the Euclidean ball with box or ellipsoidal target regions when specifications are independent or anisotropic, and introduce adaptive tolerances that reflect evolving design specifications. Integrating range-aware optimization with automated synthesis and characterization pipelines could enable closed-loop, specification-driven materials discovery.

Author contributions

S. J. contributed conceptualization, methodology, software, investigation, formal analysis, visualization, and writing of the original draft. J. W. contributed the time-dependent density functional theory calculations and data for the sequence-defined oligomer case study, and review and editing of the manuscript. C. M. S. contributed review and editing of the manuscript. M. A. W. contributed conceptualization, supervision, funding acquisition, and review and editing of the manuscript.

Conflicts of interest

There are no conflicts to declare.

Data availability

All code and data required to reproduce the results, including the benchmark datasets, the kinetic Monte Carlo polymerization inputs, and the TD-DFT-derived oligomer library, are openly available at <https://github.com/webbtheosim/target-range-optimization> and are archived on Zenodo under DOI [10.5281/zenodo.21497960](https://doi.org/10.5281/zenodo.21497960). The archived release corresponds to the version of the code used to produce the results reported here.

Supplementary information (SI) is available. See DOI: <https://doi.org/10.1039/d6dd00358c>.

Acknowledgements

S. J. and M. A. W. acknowledge support from the National Science Foundation under Grant No. 2118861. Simulations and analyses were performed using resources from Princeton Research Computing at Princeton University, which is a consortium led by the Princeton Institute for Computational Science and Engineering (PICSciE) and Research Computing in the Office of Information Technology.

References

- 1 S. R. Chitturi, A. Ramdas, Y. Wu, B. Rohr, S. Ermon, J. Dionne, F. H. D. Jornada, M. Dunne, C. Tassone, W. Neiswanger and D. Ratner, Targeted Materials Discovery Using Bayesian Algorithm Execution, *npj Comput. Mater.*, 2024, **10**, 156, DOI: [10.1038/s41524-024-01326-2](https://doi.org/10.1038/s41524-024-01326-2).
- 2 Y. Tian, T. Li, J. Pang, Y. Zhou, D. Xue, X. Ding and T. Lookman, Materials Design with Target-Oriented Bayesian Optimization, *npj Comput. Mater.*, 2025, **11**, 209, DOI: [10.1038/s41524-025-01704-4](https://doi.org/10.1038/s41524-025-01704-4).
- 3 Z. Qian, Z. Cao, L. Galuska, S. Zhang, J. Xu and X. Gu, Glass Transition Phenomenon for Conjugated Polymers, *Macromol. Chem. Phys.*, 2019, **220**, 1900062, DOI: [10.1002/macp.201900062](https://doi.org/10.1002/macp.201900062).
- 4 C. Kim, A. Chandrasekaran, A. Jha and R. Ramprasad, Active-learning and materials design: the example of high glass transition temperature polymers, *MRS Commun.*, 2019, **9**, 860–866, DOI: [10.1557/mrc.2019.78](https://doi.org/10.1557/mrc.2019.78).
- 5 L. Yu and A. Zunger, Identification of Potential Photovoltaic Absorbers Based on First-Principles Spectroscopic Screening, *Phys. Rev. Lett.*, 2012, **108**, 068701, DOI: [10.1103/PhysRevLett.108.068701](https://doi.org/10.1103/PhysRevLett.108.068701).
- 6 S. Kim, J. A. Márquez, T. Unold and A. Walsh, Upper Limit to the Photovoltaic Efficiency of Imperfect Crystals from First Principles, *Energy Environ. Sci.*, 2020, **13**, 1481–1491, DOI: [10.1039/D0EE00291G](https://doi.org/10.1039/D0EE00291G).
- 7 G. E. Eperon, T. Leijtens, K. A. Bush, R. Prasanna, T. Green, J. T.-W. Wang, D. P. McMeekin, G. Volonakis, R. L. Milot, R. May, A. Palmstrom, D. J. Slotcavage, R. A. Belisle, J. B. Patel, E. S. Parrott, R. J. Sutton, W. Ma, F. Moghadam, B. Conings, A. Babayigit, H.-G. Boyen, S. Bent, F. Giustino, L. M. Herz, M. B. Johnston, M. D. McGehee and H. J. Snaith, Perovskite-perovskite tandem photovoltaics with optimized band gaps, *Science*, 2016, **354**, 861–865, DOI: [10.1126/science.aaf9717](https://doi.org/10.1126/science.aaf9717).
- 8 B. Bhushan, *Principles and Applications of Tribology; Tribology in Practice Series*, John Wiley & Sons, 2013, 2nd edn, DOI: [10.1002/9781118403020](https://doi.org/10.1002/9781118403020).
- 9 Y.-Y. Chen and J.-H. Horng, Investigation of Lubricant Viscosity and Third-Particle Contribution to Contact Behavior in Dry and Lubricated Three-Body Contact Conditions, *Front. Mech. Eng.*, 2024, **10**, 1390335, DOI: [10.3389/fmech.2024.1390335](https://doi.org/10.3389/fmech.2024.1390335).
- 10 M. F. Ashby and D. Cebon, Materials selection in mechanical design, *J. Phys.*, 1993, **3**, 1–9, DOI: [10.1051/jp4:1993701](https://doi.org/10.1051/jp4:1993701).
- 11 T. Schofield, D. Robbins and G. Miró-Quesada, *Quality by Design for Biopharmaceutical Drug Product Development*, Springer, 2015, pp. 511–535, DOI: [10.1007/978-1-4939-2316-8_21](https://doi.org/10.1007/978-1-4939-2316-8_21).
- 12 K. T. Ulrich, S. D. Eppinger and M. C. Yang, *Product Design and Development*, 7th edn; McGraw-Hill Education, New York, NY, 2020.
- 13 M. Ashby, Y. Bréchet, D. Cebon and L. Salvo, Selection strategies for materials and processes, *Mater. Des.*, 2004, **25**, 51–67, DOI: [10.1016/S0261-3069\(03\)00159-6](https://doi.org/10.1016/S0261-3069(03)00159-6).
- 14 C. Zeni, R. Pinsler, D. Zügner, A. Fowler, M. Horton, X. Fu, Z. Wang, A. Shysheya, J. Crabbé, S. Ueda, R. Sordillo, L. Sun, J. Smith, B. Nguyen, H. Schulz, S. Lewis, C.-W. Huang, Z. Lu, Y. Zhou, H. Yang, H. Hao, J. Li, C. Yang, W. Li, R. Tomioka and T. Xie, A generative model for inorganic materials design, *Nature*, 2025, **639**, 624–632, DOI: [10.1038/s41586-025-08628-5](https://doi.org/10.1038/s41586-025-08628-5).
- 15 P. Renz, S. Luukkonen and G. Klambauer, Diverse hits in *de novo* molecule design: Diversity-based comparison of goal-directed generators, *J. Chem. Inf. Model.*, 2024, **64**, 5756–5761, DOI: [10.1021/acs.jcim.4c00519](https://doi.org/10.1021/acs.jcim.4c00519).
- 16 A. S. Alshehri, A. K. Tula, F. You and R. Gani, Next generation pure component property estimation models: With and without machine learning techniques, *AIChE J.*, 2022, **68**, e17469, DOI: [10.1002/aic.17469](https://doi.org/10.1002/aic.17469).
- 17 T. H. Kim, S. Pahari and J. S.-I. Kwon, SigmaFormer: Augmenting transformer encoders with COSMO sigma profiles for pure component property prediction, *AIChE J.*, 2026, e70450, DOI: [10.1002/aic.70450](https://doi.org/10.1002/aic.70450).
- 18 S. Pahari, C. H. Lee, N. Sitapure and J. S.-I. Kwon, Predicting both thermodynamic and kinetic properties of crystallizing molecules *via* transformer-based language model, *AIChE J.*, 2025, **71**, e18923, DOI: [10.1002/aic.18923](https://doi.org/10.1002/aic.18923).
- 19 J. Moćkus, On Bayesian Methods for Seeking the Extremum, *Optimization Techniques IFIP Technical Conference Novosibirsk, July 1–7, Berlin, Heidelberg, 1975*, pp. 400–404, DOI: [10.1007/3-540-07165-2_55](https://doi.org/10.1007/3-540-07165-2_55).
- 20 D. R. Jones, M. Schonlau and W. J. Welch, Efficient Global Optimization of Expensive Black-Box Functions, *J. Global Optim.*, 1998, **13**, 455–492, DOI: [10.1023/A:1008306431147](https://doi.org/10.1023/A:1008306431147).
- 21 J. Snoek, H. Larochelle and R. P. Adams, Practical Bayesian Optimization of Machine Learning Algorithms, *Advances in Neural Information Processing Systems*, 2012, **25**, pp. 2951–2959, DOI: [10.5555/2999325.2999464](https://doi.org/10.5555/2999325.2999464).
- 22 B. J. Shields, J. Stevens, J. Li, M. Parasram, F. Damani, J. I. M. Alvarado, J. M. Janey, R. P. Adams and A. G. Doyle, Bayesian Reaction Optimization as a Tool for Chemical Synthesis, *Nature*, 2021, **590**, 89–96, DOI: [10.1038/s41586-021-03213-y](https://doi.org/10.1038/s41586-021-03213-y).
- 23 C. J. Taylor, K. C. Felton, D. Wigh, M. I. Jeraal, R. Grainger, G. Chessari, C. N. Johnson and A. A. Lapkin, Accelerated Chemical Reaction Optimization Using Multi-Task Learning, *ACS Cent. Sci.*, 2023, **9**, 957–968, DOI: [10.1021/acscentsci.3c00050](https://doi.org/10.1021/acscentsci.3c00050).
- 24 J. K. Pedersen, C. M. Clausen, O. A. Krysiak, B. Xiao, T. A. A. Batchelor, T. Löffler, V. A. Mints, L. Banko, M. Arenz, A. Savan, W. Schuhmann, A. Ludwig and J. Rossmeisl, Bayesian Optimization of High-Entropy Alloy Compositions for Electrocatalytic Oxygen Reduction, *Angew. Chem., Int. Ed.*, 2021, **60**, 24144–24152, DOI: [10.1002/anie.202108116](https://doi.org/10.1002/anie.202108116).
- 25 D. Khatamsaz, B. Vela, P. Singh, D. D. Johnson, D. Allaire and R. Arróyave, Bayesian Optimization with Active Learning of Design Constraints Using an Entropy-Based Approach, *npj Comput. Mater.*, 2023, **9**, 49, DOI: [10.1038/s41524-023-01006-7](https://doi.org/10.1038/s41524-023-01006-7).

- 26 Y. An, M. A. Webb and W. M. Jacobs, Active learning of the thermodynamics-dynamics trade-off in protein condensates, *Sci. Adv.*, 2024, **10**, eadj2448, DOI: [10.1126/sciadv.adj2448](https://doi.org/10.1126/sciadv.adj2448).
- 27 C. Y. P. Wilding, R. A. Bourne and N. J. Warren, Integrating Mechanistic Modelling with Bayesian Optimisation: Accelerated Self-Driving Laboratories for RAFT Polymerisation, *Digital Discovery*, 2025, **4**, 2797–2803, DOI: [10.1039/D5DD000258C](https://doi.org/10.1039/D5DD000258C).
- 28 R. J. Dalal, F. Oviedo, M. C. Leyden and T. M. Reineke, Polymer Design via SHAP and Bayesian Machine Learning Optimizes pDNA and CRISPR Ribonucleoprotein Delivery, *Chem. Sci.*, 2024, **15**, 7219–7228, DOI: [10.1039/D3SC06920F](https://doi.org/10.1039/D3SC06920F).
- 29 S. Jiang and M. A. Webb, Generative Active Learning across Polymer Architectures and Solvophobicities for Targeted Rheological Behavior, *npj Comput. Mater.*, 2026, **12**, 28, DOI: [10.1038/s41524-025-01900-2](https://doi.org/10.1038/s41524-025-01900-2).
- 30 H. J. Kushner, A New Method of Locating the Maximum Point of an Arbitrary Multipeak Curve in the Presence of Noise, *J. Basic Eng.*, 1964, **86**, 97–106, DOI: [10.1115/1.3653121](https://doi.org/10.1115/1.3653121).
- 31 N. Srinivas, A. Krause, S. M. Kakade and M. Seeger, Gaussian Process Optimization in the Bandit Setting: No Regret and Experimental Design, *Proceedings of the 27th International Conference on International Conference on Machine Learning*, 2010, pp. 1015–1022, DOI: [10.5555/3104322.3104451](https://doi.org/10.5555/3104322.3104451).
- 32 J. Knowles, ParEGO: A Hybrid Algorithm With On-Line Landscape Approximation for Expensive Multiobjective Optimization Problems, *IEEE Transactions on Evolutionary Computation*, 2006, vol. 10, pp. 50–66, DOI: [10.1109/TEVC.2005.851274](https://doi.org/10.1109/TEVC.2005.851274).
- 33 M. T. M. Emmerich, A. H. Deutz and J. W. Klinkenberg, Hypervolume-Based Expected Improvement: Monotonicity Properties and Exact Computation, *IEEE Congress on Evolutionary Computation (CEC)*, 2011, pp. 2147–2154, DOI: [10.1109/CEC.2011.5949880](https://doi.org/10.1109/CEC.2011.5949880).
- 34 S. Daulton, M. Balandat and E. Bakshy, Differentiable Expected Hypervolume Improvement for Parallel Multi-Objective Bayesian Optimization, *Advances in Neural Information Processing Systems*, 2020.
- 35 A. Gotovos, N. Casati, G. Hitz and A. Krause, Active Learning for Level Set Estimation. *Proceedings of the Twenty-Third International Joint Conference on Artificial Intelligence (IJCAI 2013)*, 2013, pp. 1344–1350.
- 36 B. Bryan, J. G. Schneider, R. C. Nichol, C. J. Miller, C. R. Genovese and L. A. Wasserman, Active Learning for Identifying Function Threshold Boundaries, *Advances in Neural Information Processing Systems*, NeurIPS 2005, 2005, vol. 18, pp. 163–170.
- 37 A. N. Marques, R. R. Lam and K. E. Willcox, Contour Location via Entropy Reduction Leveraging Multiple Information Sources, *Advances in Neural Information Processing Systems (NeurIPS 2018)*, 2018, DOI: [10.5555/3327345.3327428](https://doi.org/10.5555/3327345.3327428).
- 38 B. J. Bichon, M. S. Eldred, L. P. Swiler, S. Mahadevan and J. M. McFarland, Efficient global reliability analysis for nonlinear implicit performance functions, *AIAA J.*, 2008, **46**, 2459–2468, DOI: [10.2514/1.34321](https://doi.org/10.2514/1.34321).
- 39 W. Neiswanger, K. A. Wang and S. Ermon, Bayesian Algorithm Execution: Estimating Computable Properties of Black-Box Functions Using Mutual Information, *Proceedings of the 38th International Conference on Machine Learning*, 2021, pp. 8005–8015.
- 40 A. Nakao, H. Kaneko and K. Funatsu, Development of an Adaptive Experimental Design Method Based on Probability of Achieving a Target Range through Parallel Experiments, *Ind. Eng. Chem. Res.*, 2016, **55**, 5726–5735, DOI: [10.1021/acs.iecr.6b00852](https://doi.org/10.1021/acs.iecr.6b00852).
- 41 J. R. Gardner, M. J. Kusner, Z. E. Xu, K. Q. Weinberger and J. P. Cunningham, Bayesian Optimization with Inequality Constraints, *Proceedings of the 31st International Conference on Machine Learning*, 2014, pp. 937–945.
- 42 M. A. Gelbart, J. Snoek and R. P. Adams, Bayesian Optimization with Unknown Constraints, *Proceedings of the Thirtieth Conference on Uncertainty in Artificial Intelligence*, 2014, pp. 250–259.
- 43 D. Eriksson and M. Poloczek, Scalable Constrained Bayesian Optimization, *Proceedings of the 24th International Conference on Artificial Intelligence and Statistics*, 2021, vol. 130, pp. 730–738.
- 44 S. Ariafar, J. Coll-Font, D. H. Brooks and J. G. Dy, ADMMBO: Bayesian Optimization with Unknown Constraints Using ADMM, *J. Mach. Learn. Res.*, 2019, **20**, 1–26.
- 45 C. Lu and J. A. Paulson, No-Regret Bayesian Optimization with Unknown Equality and Inequality Constraints using Exact Penalty Functions, *IFAC-PapersOnLine*, 2022, vol. 55, pp. 895–902, DOI: [10.1016/j.ifacol.2022.07.558](https://doi.org/10.1016/j.ifacol.2022.07.558).
- 46 A. Kudva, F. Sorourifar and J. A. Paulson, Constrained robust Bayesian optimization of expensive noisy black-box functions with guaranteed regret bounds, *AIChe J.*, 2022, **68**, e17857, DOI: [10.1002/aic.17857](https://doi.org/10.1002/aic.17857).
- 47 P. Jaiswal, H. Honnappa and V. A. Rao, Bayesian Joint Chance Constrained Optimization: Approximations and Statistical Consistency, *SIAM J. Optim.*, 2023, **33**, 1968–1995, DOI: [10.1137/21M1430005](https://doi.org/10.1137/21M1430005).
- 48 J. K. Pugh, L. B. Soros and K. O. Stanley, Quality Diversity: A New Frontier for Evolutionary Computation, *Front. Robot. AI*, 2016, **3**, 40, DOI: [10.3389/frobt.2016.00040](https://doi.org/10.3389/frobt.2016.00040).
- 49 M. Konakovic Lukovic, Y. Tian and W. Matusik, Diversity-guided multi-objective Bayesian optimization with batch evaluations, *Advances in Neural Information Processing Systems*, 2020, vol. 33, pp. 17708–17720.
- 50 N. Maus, K. Wu, D. Eriksson and J. Gardner, Discovering Many Diverse Solutions with Bayesian Optimization, *arXiv*, 2022, preprint, arXiv:2210.10953, DOI: [10.48550/arXiv.2210.10953](https://doi.org/10.48550/arXiv.2210.10953).
- 51 F. Häse, L. M. Roch and A. Aspuru-Guzik, Chimera: enabling hierarchy based multi-objective optimization for self-driving laboratories, *Chem. Sci.*, 2018, **9**, 7642–7655, DOI: [10.1039/C8SC02239A](https://doi.org/10.1039/C8SC02239A).
- 52 P. M. Attia, A. Grover, N. Jin, K. A. Severson, T. M. Markov, Y.-H. Liao, M. H. Chen, B. Cheong, N. Perkins, Z. Yang, P. K. Herring, M. Aykol, S. J. Harris, R. D. Braatz,

- S. Ermon and W. C. Chueh, Closed-loop optimization of fast-charging protocols for batteries with machine learning, *Nature*, 2020, **578**, 397–402, DOI: [10.1038/s41586-020-1994-5](https://doi.org/10.1038/s41586-020-1994-5).
- 53 A. G. Kusne, H. Yu, C. Wu, H. Zhang, J. Hattrick-Simpers, B. DeCost, S. Sarker, C. Oses, C. Toher, S. Curtarolo, A. V. Davydov, R. Agarwal, L. A. Bendersky, M. Li, A. Mehta and I. Takeuchi, On-the-fly closed-loop materials discovery via Bayesian active learning, *Nat. Commun.*, 2020, **11**, 5966, DOI: [10.1038/s41467-020-19597-w](https://doi.org/10.1038/s41467-020-19597-w).
- 54 N. J. Szymanski, B. Rendy, Y. Fei, R. E. Kumar, T. He, D. Milsted, M. J. McDermott, M. Gallant, E. D. Cubuk, A. Merchant, H. Kim, A. Jain, C. J. Bartel, K. Persson, Y. Zeng and G. Ceder, An autonomous laboratory for the accelerated synthesis of inorganic materials, *Nature*, 2023, **624**, 86–91, DOI: [10.1038/s41586-023-06734-w](https://doi.org/10.1038/s41586-023-06734-w).
- 55 F. Ren, L. Ward, T. Williams, K. J. Laws, C. Wolverton, J. Hattrick-Simpers and A. Mehta, Accelerated discovery of metallic glasses through iteration of machine learning and high-throughput experiments, *Sci. Adv.*, 2018, **4**, eaq1566, DOI: [10.1126/sciadv.aq1566](https://doi.org/10.1126/sciadv.aq1566).
- 56 M. J. Tamasi, R. A. Patel, C. H. Borca, S. Kosuri, H. Mugnier, R. Upadhyaya, N. S. Murthy, M. A. Webb and A. J. Gormley, Machine learning on a robotic platform for the design of polymer–protein hybrids, *Adv. Mater.*, 2022, **34**, 2201809, DOI: [10.1002/adma.202201809](https://doi.org/10.1002/adma.202201809).
- 57 C. E. Rasmussen and C. K. I. Williams, *Gaussian Processes for Machine Learning*, MIT Press, Cambridge, MA, 2006, DOI: [10.7551/mitpress/3206.001.0001](https://doi.org/10.7551/mitpress/3206.001.0001).
- 58 M. Aldeghi, D. E. Graff, N. Frey, J. A. Morrone, E. O. Pyzer-Knapp, K. E. Jordan and C. W. Coley, Roughness of Molecular Property Landscapes and Its Impact on Modellability, *J. Chem. Inf. Model.*, 2022, **62**, 4660–4671, DOI: [10.1021/acs.jcim.2c00903](https://doi.org/10.1021/acs.jcim.2c00903).
- 59 M. D. McKay, R. J. Beckman and W. J. Conover, A Comparison of Three Methods for Selecting Values of Input Variables in the Analysis of Output from a Computer Code, *Technometrics*, 1979, **21**, 239–245, DOI: [10.1080/00401706.1979.10489755](https://doi.org/10.1080/00401706.1979.10489755).
- 60 T. Desautels, A. Krause and J. W. Burdick, Parallelizing Exploration-Exploitation Tradeoffs in Gaussian Process Bandit Optimization, *J. Mach. Learn. Res.*, 2014, **15**, 4053–4103.
- 61 C. X. Cheng, R. Astudillo, T. Desautels and Y. Yue, Practical Bayesian Algorithm Execution via Posterior Sampling, *Advances in Neural Information Processing Systems*, 2024, vol. 37, pp. 135186–135207, DOI: [10.52202/079017-4296](https://doi.org/10.52202/079017-4296).
- 62 A. K. Uhrenholt and B. S. Jensen, Efficient Bayesian Optimization for Target Vector Estimation–, *Proceedings of the Twenty-Second International Conference on Artificial Intelligence and Statistics*, 2019, vol. 89, pp. 2661–2670.
- 63 R.-R. Griffiths and J. M. Hernández-Lobato, Constrained Bayesian optimization for automatic chemical design using variational autoencoders, *Chem. Sci.*, 2020, **11**, 577–586, DOI: [10.1039/C9SC04026A](https://doi.org/10.1039/C9SC04026A).
- 64 A. Rahimi and B. Recht, Random Features for Large-Scale Kernel Machines, *Advances in Neural Information Processing Systems*, 2007, pp. 1177–1184.
- 65 J. T. Wilson, V. Borovitskiy, A. Terenin, P. Mostowsky and M. P. Deisenroth, Pathwise Conditioning of Gaussian Processes, *J. Mach. Learn. Res.*, 2021, **22**, 1–47.
- 66 K. Kandasamy, A. Krishnamurthy, J. Schneider and B. Póczos, Parallelised Bayesian Optimisation via Thompson Sampling, *Proceedings of the 21st International Conference on Artificial Intelligence and Statistics*, 2018, pp. 133–142.
- 67 M. Jamil and X.-S. Yang, A literature survey of benchmark functions for global optimisation problems, *Int. J. Math. Model. Numer. Optimisation*, 2013, **4**, 150–194, DOI: [10.1504/IJMMNO.2013.055204](https://doi.org/10.1504/IJMMNO.2013.055204).
- 68 F. Pellegrino, R. Isopescu, L. Pellutì, F. Sordello, A. M. Rossi, E. Ortel, G. Martra, V.-D. Hodoroaba and V. Maurino, Machine Learning Approach for Elucidating and Predicting the Role of Synthesis Parameters on the Shape and Size of TiO₂ Nanoparticles, *Sci. Rep.*, 2020, **10**, 18910, DOI: [10.1038/s41598-020-75967-w](https://doi.org/10.1038/s41598-020-75967-w).
- 69 W. W. Oliver, W. M. Jacobs and M. A. Webb, When B2 is Not Enough: Evaluating Simple Metrics for Predicting Phase Separation of Intrinsically Disordered Proteins, *J. Phys. Chem. B*, 2025, **129**, 9551–9565, DOI: [10.1021/acs.jpcb.5c04955](https://doi.org/10.1021/acs.jpcb.5c04955).
- 70 S. Jiang, A. B. Dieng and M. A. Webb, Property-guided generation of complex polymer topologies using variational autoencoders, *npj Comput. Mater.*, 2024, **10**, 139, DOI: [10.1038/s41524-024-01328-0](https://doi.org/10.1038/s41524-024-01328-0).
- 71 D. L. Mobley and J. P. Guthrie, FreeSolv: A Database of Experimental and Calculated Hydration Free Energies, with Input Files, *J. Comput.-Aided Mol. Des.*, 2014, **28**, 711–720, DOI: [10.1007/s10822-014-9747-x](https://doi.org/10.1007/s10822-014-9747-x).
- 72 J. S. Delaney, ESOL: Estimating Aqueous Solubility Directly from Molecular Structure, *J. Chem. Inf. Comput. Sci.*, 2004, **44**, 1000–1005, DOI: [10.1021/ci034243x](https://doi.org/10.1021/ci034243x).
- 73 R. Ramakrishnan, P. O. Dral, M. Rupp and O. A. von Lilienfeld, Quantum Chemistry Structures and Properties of 134 Kilo Molecules, *Sci. Data*, 2014, **1**, 140022, DOI: [10.1038/sdata.2014.22](https://doi.org/10.1038/sdata.2014.22).
- 74 Q. M. Gallagher and M. A. Webb, Data Efficiency of Classification Strategies for Chemical and Materials Design, *Digital Discovery*, 2025, **4**, 135–148, DOI: [10.1039/D4DD00298A](https://doi.org/10.1039/D4DD00298A).
- 75 Z. Wu, B. Ramsundar, E. N. Feinberg, J. Gomes, C. Geniesse, A. S. Pappu, K. Leswing and V. Pande, MoleculeNet: a benchmark for molecular machine learning, *Chem. Sci.*, 2018, **9**, 513–530, DOI: [10.1039/C7SC02664A](https://doi.org/10.1039/C7SC02664A).
- 76 M. K. Lenzi, M. F. Cunningham, E. L. Lima and J. C. Pinto, Producing bimodal molecular weight distribution polymer resins using living and conventional free-radical polymerization, *Ind. Eng. Chem. Res.*, 2005, **44**, 2568–2578, DOI: [10.1021/ie0496479](https://doi.org/10.1021/ie0496479).
- 77 S. K. Fierens, D. R. D'hooge, P. H. M. Van Steenberge, M.-F. Reyniers and G. B. Marin, MAMA-SG1 Initiated Nitroxide Mediated Polymerization of Styrene: From Arrhenius Parameters to Model-Based Design, *Chem. Eng. J.*, 2015, **278**, 407–420, DOI: [10.1016/j.cej.2014.09.024](https://doi.org/10.1016/j.cej.2014.09.024).

- 78 D. T. Gillespie, A General Method for Numerically Simulating the Stochastic Time Evolution of Coupled Chemical Reactions, *J. Comput. Phys.*, 1976, **22**, 403–434, DOI: [10.1016/0021-9991\(76\)90041-3](https://doi.org/10.1016/0021-9991(76)90041-3).
- 79 D. T. Gillespie, Exact Stochastic Simulation of Coupled Chemical Reactions, *J. Phys. Chem.*, 1977, **81**, 2340–2361, DOI: [10.1021/j100540a008](https://doi.org/10.1021/j100540a008).
- 80 T. Romero Pietrafesa, A. D. Trigilio, Y. W. Marien, P. Reyes, M. Edeleva, M. Asteasuain, P. H. Van Steenberge and D. R. D'hooge, Benchmark cases and guidelines for kinetic Monte Carlo simulations with linear polymers, *Ind. Eng. Chem. Res.*, 2025, **64**, 9974–9992, DOI: [10.1021/acs.iecr.5c00639](https://doi.org/10.1021/acs.iecr.5c00639).
- 81 E. Runge and E. K. U. Gross, Density-Functional Theory for Time-Dependent Systems, *Phys. Rev. Lett.*, 1984, **52**, 997–1000, DOI: [10.1103/PhysRevLett.52.997](https://doi.org/10.1103/PhysRevLett.52.997).
- 82 M. E. Casida, Recent Advances in Density Functional Methods, *World Sci.*, 1995, **1**, 155–192.
- 83 N. Nozaki, A. Uva, T. Iwahashi, H. Matsumoto, H. Tran and M. Ashizawa, Impact of aromatic to quinoidal transformation on the degradation kinetics of imine-based semiconducting polymers, *RSC Appl. Polym.*, 2025, **3**, 257–267, DOI: [10.1039/D4LP00310A](https://doi.org/10.1039/D4LP00310A).
- 84 RDKit, *RDKit: Open-source cheminformatics*, <https://www.rdkit.org>, 2006, DOI: [10.5281/zenodo.591637](https://doi.org/10.5281/zenodo.591637).
- 85 S. Wang, J. Witek, G. A. Landrum and S. Riniker, Improving Conformer Generation for Small Rings and Macrocycles Based on Distance Geometry and Experimental Torsional-Angle Preferences, *J. Chem. Inf. Model.*, 2020, **60**, 2044–2058, DOI: [10.1021/acs.jcim.0c00025](https://doi.org/10.1021/acs.jcim.0c00025).
- 86 T. A. Halgren, Merck molecular force field. I. Basis, form, scope, parameterization, and performance of MMFF94, *J. Comput. Chem.*, 1996, **17**, 490–519, DOI: [10.1002/\(SICI\)1096-987X\(199604\)17:5<490::AID-JCC1>3.0.CO;2-P](https://doi.org/10.1002/(SICI)1096-987X(199604)17:5<490::AID-JCC1>3.0.CO;2-P).
- 87 R. B. Roseli, P. C. Tapping and T. W. Kee, Origin of the excited-state absorption spectrum of polythiophene, *J. Phys. Chem. Lett.*, 2017, **8**, 2806–2811, DOI: [10.1021/acs.jpclett.7b01053](https://doi.org/10.1021/acs.jpclett.7b01053).
- 88 F. Neese, F. Wennmohs, U. Becker and C. Riplinger, The ORCA quantum chemistry program package, *J. Chem. Phys.*, 2020, **152**, 224108, DOI: [10.1063/5.0004608](https://doi.org/10.1063/5.0004608).
- 89 C. Bannwarth, S. Ehlert and S. Grimme, GFN2-xTB—An Accurate and Broadly Parametrized Self-Consistent Tight-Binding Quantum Chemical Method with Multipole Electrostatics and Density-Dependent Dispersion Contributions, *J. Chem. Theory Comput.*, 2019, **15**, 1652–1671, DOI: [10.1021/acs.jctc.8b01176](https://doi.org/10.1021/acs.jctc.8b01176).
- 90 S. Ehlert, M. Stahn, S. Spicher and S. Grimme, Robust and Efficient Implicit Solvation Model for Fast Semiempirical Methods, *J. Chem. Theory Comput.*, 2021, **17**, 4250–4261, DOI: [10.1021/acs.jctc.1c00471](https://doi.org/10.1021/acs.jctc.1c00471).
- 91 M. Müller, A. Hansen and S. Grimme, ω B97X-3c: A composite range-separated hybrid DFT method with a molecule-optimized polarized valence double- ζ basis set, *J. Chem. Phys.*, 2023, **158**, 014103, DOI: [10.1063/5.0133026](https://doi.org/10.1063/5.0133026).
- 92 V. Barone and M. Cossi, Quantum Calculation of Molecular Energies and Energy Gradients in Solution by a Conductor Solvent Model, *J. Phys. Chem. A*, 1998, **102**, 1995–2001, DOI: [10.1021/jp9714387](https://doi.org/10.1021/jp9714387).
- 93 H. Moriwaki, Y.-S. Tian, N. Kawashita and T. Takagi, Mordred: a molecular descriptor calculator, *J. Cheminf.*, 2018, **10**, 4, DOI: [10.1186/s13321-018-0258-y](https://doi.org/10.1186/s13321-018-0258-y).
- 94 I. T. Jolliffe, *Principal Component Analysis*, Springer Series in Statistics, Springer: New York, 2nd ed., 2002, DOI: [10.1007/b98835](https://doi.org/10.1007/b98835).
- 95 L. Kaufman and P. J. Rousseeuw, *Finding Groups in Data: An Introduction to Cluster Analysis*, John Wiley & Sons, New York, 1990, DOI: [10.1002/9780470316801](https://doi.org/10.1002/9780470316801).
- 96 J. MacQueen, Some Methods for Classification and Analysis of Multivariate Observations. *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability Statistics*, vol. 1, Berkeley, CA, 1967, pp. 281–297.
- 97 D. R. Jones, A Taxonomy of Global Optimization Methods Based on Response Surfaces, *J. Global Optim.*, 2001, **21**, 345–383, DOI: [10.1023/a:1012771025575](https://doi.org/10.1023/a:1012771025575).
- 98 D. T. Gentekos, L. N. Dupuis and B. P. Fors, Beyond dispersity: Deterministic control of polymer molecular weight distribution, *J. Am. Chem. Soc.*, 2016, **138**, 1848–1851, DOI: [10.1021/jacs.5b13565](https://doi.org/10.1021/jacs.5b13565).
- 99 D. T. Gentekos, R. M. Sifri and B. P. Fors, Controlling polymer properties through the shape of the molecular-weight distribution, *Nat. Rev. Mater.*, 2019, **4**, 761–774, DOI: [10.1038/s41578-019-0138-8](https://doi.org/10.1038/s41578-019-0138-8).
- 100 K.-Y. Yoo, B.-G. Jeong and H.-K. Rhee, Molecular Weight Distribution Control in a Batch Polymerization Reactor Using the On-Line Two-Step Method, *Ind. Eng. Chem. Res.*, 1999, **38**, 4805–4814, DOI: [10.1021/ie980799b](https://doi.org/10.1021/ie980799b).
- 101 H. Li, C. R. Collins, T. G. Ribelli, K. Matyjaszewski, G. J. Gordon, T. Kowalewski and D. J. Yaron, Tuning the molecular weight distribution from atom transfer radical polymerization using deep reinforcement learning, *Mol. Syst. Des. Eng.*, 2018, **3**, 496–508, DOI: [10.1039/C7ME00131B](https://doi.org/10.1039/C7ME00131B).
- 102 D. J. Walsh, D. A. Schinski, R. A. Schneider and D. Guironnet, General route to design polymer molecular weight distributions through flow chemistry, *Nat. Commun.*, 2020, **11**, 3094, DOI: [10.1038/s41467-020-16874-6](https://doi.org/10.1038/s41467-020-16874-6).
- 103 K. Liu, N. Corrigan, A. Postma, G. Moad and C. Boyer, A Comprehensive Platform for the Design and Synthesis of Polymer Molecular Weight Distributions, *Macromolecules*, 2020, **53**, 8867–8882, DOI: [10.1021/acs.macromol.0c01954](https://doi.org/10.1021/acs.macromol.0c01954).
- 104 H. Liu, Y.-H. Xue, Y.-L. Zhu, F.-L. Gu and Z.-Y. Lu, Inverse Design of Molecular Weight Distribution in Controlled Polymerization via a One-Pot Reaction Strategy, *Macromolecules*, 2020, **53**, 6409–6419, DOI: [10.1021/acs.macromol.0c01383](https://doi.org/10.1021/acs.macromol.0c01383).
- 105 A. MacKenzie, J. Schneider, J. Meyer and C. Loschen, Computer aided recipe design: optimization of polydisperse chemical mixtures using molecular descriptors, *React. Chem. Eng.*, 2024, **9**, 1061–1076, DOI: [10.1039/D3RE00601H](https://doi.org/10.1039/D3RE00601H).

- 106 H. Zhou, Y. Fang and H. Gao, Using Active Learning for the Computational Design of Polymer Molecular Weight Distributions, *ACS Eng. Au*, 2024, **4**, 231–240, DOI: [10.1021/acsengineeringau.3c00056](https://doi.org/10.1021/acsengineeringau.3c00056).
- 107 J. Fiosina, P. Sievers, M. Drache and S. Beuermann, Machine learning supported evolutionary optimization for multi-objective reverse engineering of radical polymerizations, *Comput. Chem. Eng.*, 2025, **199**, 109125, DOI: [10.1016/j.compchemeng.2025.109125](https://doi.org/10.1016/j.compchemeng.2025.109125).
- 108 Y. Fang, Y. Jin and H. Gao, A unified kinetic Monte Carlo and Bayesian optimization framework for the circular design of poly(methyl methacrylate): From synthesis to recycling, *Chem. Eng. J.*, 2026, **538**, 176625, DOI: [10.1016/j.cej.2026.176625](https://doi.org/10.1016/j.cej.2026.176625).
- 109 L. McInnes, J. Healy, N. Saul and L. Grossberger, UMAP: Uniform Manifold Approximation and Projection, *J. Open Source Softw.*, 2018, **3**, 861, DOI: [10.21105/joss.00861](https://doi.org/10.21105/joss.00861).