



Showcasing research from Professor Xiaonan Wang and the Intelligent Chemical Engineering Research Center, Department of Chemical Engineering at Tsinghua University, in collaboration with global leaders in AI-for-Science research.

A unified active learning framework for photosensitizer design

This work introduces a general and scalable active learning framework that streamlines molecular design workflows and substantially enhances exploration efficiency for chemists. By integrating machine learning calibrated computations, a large molecular dataset, and adaptive data acquisition, the framework accelerates the identification of promising candidates while reducing computational cost. The study demonstrates a broadly applicable approach for navigating vast chemical spaces and enabling more efficient, data-driven discovery in photochemical materials research.

Image reproduced by permission of Xiaonan Wang from *Chem. Sci.*, 2026, **17**, 916. Image enhanced with AI (Nanobanana).

As featured in:



See Jiali Li, Rafael Gómez-Bombarelli, Aron Walsh, Xiaonan Wang *et al.*, *Chem. Sci.*, 2026, **17**, 916.

Cite this: *Chem. Sci.*, 2026, 17, 916

All publication charges for this article have been paid for by the Royal Society of Chemistry

A unified active learning framework for photosensitizer design

Yizhe Chen,^{†a} Shomik Verma,^{†b} Kevin P. Greenman,^{†cde} Haoyu Yin,^a Zhihao Wang,^{id a} Lanjing Wang,^f Jiali Li,^{*g} Rafael Gómez-Bombarelli,^{id *h} Aron Walsh^{id *i} and Xiaonan Wang^{id *a}

The design of high-performance photosensitizers for next-generation photovoltaic and clean energy applications remains a formidable challenge due to the vast chemical space, competing photophysical trade-offs, and computational limitations of traditional quantum chemistry methods. While machine learning offers potential solutions, existing approaches suffer from data scarcity and inefficient exploration of molecular configurations. This work introduces a unified active learning framework that systematically integrates semi-empirical quantum calculations with adaptive molecular screening strategies to accelerate photosensitizer discovery. Our methodology combines three principal components: (1) A hybrid quantum mechanics/machine learning pipeline generating a chemically diverse molecular dataset while maintaining quantum chemical accuracy at significantly reduced computational costs; (2) a graph neural network architecture and uncertainty quantification; (3) Novel acquisition strategies that dynamically balance broad chemical space exploration with targeted optimization of photophysical objectives. The framework demonstrates superior performance in predicting critical energy levels (T_1/S_1) compared to conventional screening approaches, while effectively prioritizing synthetically feasible candidates. By open-sourcing both the curated molecular dataset and implementation tools, this work establishes an extensible platform for data-driven discovery of optoelectronic materials, with immediate applications in solar energy conversion and beyond.

Received 30th July 2025

Accepted 9th November 2025

DOI: 10.1039/d5sc05749c

rsc.li/chemical-science

1 Introduction

Photosensitizers (PSs) have emerged as critical functional materials in modern energy and biomedical technologies, driving innovations from solar energy harvesting to

photodynamic therapy.^{1–4} Their expanding applications in wearable devices, sterilization systems, and optical sensors demand precise control over photophysical properties such as triplet yields and excited-state lifetimes.^{5,6} However, the rational design of high-performance PSs faces three fundamental challenges that hinder rapid progress.

First, previous work has shown that combinatorial D–A assembly yields libraries containing more than one million PS candidates, far exceeding the capacity of conventional trial-and-error approaches.^{7–9} For example, subtle structural variations in porphyrin derivatives—such as peripheral substituent patterns—can shift absorption maxima by over 50 nm while altering quantum yields by orders of magnitude.^{7,10,11} Second, the intricate balance between competing photophysical properties creates a complex optimization landscape. A PS optimized for strong visible light absorption may suffer from rapid triplet–triplet annihilation, while molecules with ideal S_1/T_1 energy ratios often exhibit poor solubility or photostability.^{12,13} Third, computational screening methods such as time-dependent density-functional theory (TD-DFT), though theoretically rigorous, become prohibitively expensive for geometry optimization and large-scale exploration, requiring days of computation for a single medium-sized molecule (50+ atoms).^{14,15}

^aDepartment of Chemical Engineering, State Key Laboratory of Chemical Engineering and Low-carbon Technology, Tsinghua University, Beijing 100084, China. E-mail: wangxiaonan@tsinghua.edu.cn

^bDepartment of Mechanical Engineering, Massachusetts Institute of Technology, Cambridge, MA 02139, USA. E-mail: skverma@mit.edu

^cDepartment of Chemical Engineering, Catholic Institute of Technology, Cambridge, Massachusetts, USA. E-mail: kgreenman@catholic.tech

^dDepartment of Chemistry, Catholic Institute of Technology, Cambridge, Massachusetts, USA

^eDepartment of Chemical Engineering, Massachusetts Institute of Technology, Cambridge, Massachusetts, USA

^fInstitute of Flexible Electronics (IFE) & Frontiers Science Center for Flexible Electronics, Northwestern Polytechnical University, Xi'an, Shaanxi 710072, China. E-mail: wanglanjing@mail.nwpu.edu.cn

^gResearch Center for Eco-Environmental Sciences, Chinese Academy of Sciences, Beijing 100085, China. E-mail: jllli@rcees.ac.cn

^hDepartment of Materials Science and Engineering, Massachusetts Institute of Technology, Cambridge, MA 02139, USA. E-mail: rafagb@mit.edu

ⁱDepartment of Materials, Imperial College London, London SW7 2AZ, UK. E-mail: a.walsh@imperial.ac.uk

[†] These authors contributed equally to this work.

Recent advances in machine learning (ML) offer promising solutions to these bottlenecks.^{16,17} By establishing quantitative structure–property relationships (QSPRs), ML models can predict key PS characteristics like singlet-triplet energy gaps (ΔE_{ST}) with millisecond inference times.¹⁰ Nevertheless, existing ML approaches face two critical limitations:¹ Public datasets contain less than 0.1% of the required photophysical data for PS design, creating severe data scarcity issues;² conventional ML workflows prioritize passive learning from static datasets, inefficiently allocating computational resources to chemically redundant regions.^{12–14}

Active learning (AL) is a machine learning approach where the model selects the most informative data points for labeling, aiming to improve performance with fewer labeled examples.¹⁸ It addresses these limitations through iterative cycles of prediction and targeted data acquisition.^{15,19–21} Unlike traditional methods that treat all molecules equally, AL algorithms dynamically identify the most informative candidates for quantum chemical calculations—those with high prediction uncertainty or high potential to improve model performance.^{22–24} Recent demonstrations in catalyst discovery achieved $32\times$ acceleration over random screening by prioritizing metal alloys with optimal d-band centers.^{15,25} For PS design specifically, AL's ability to navigate high-dimensional chemical spaces while respecting synthetic constraints could revolutionize molecular discovery pipelines.^{26–30}

Despite these promising developments, existing AL applications in molecular and materials discovery remain limited in generality and methodological scope. Yoo *et al.* developed an iterative inverse-design workflow combining density-functional tight binding (DFTB) for property labeling, a graph convolutional neural network (GCNN) surrogate, and a masked-language-model (MLM) generator to optimize the HOMO–LUMO gap (HLG) of molecules.³¹ Their selection strategy primarily relies on property-value thresholds and periodic surrogate retraining, rather than a formal uncertainty-based acquisition loop. In addition, the target property (HLG) is not directly related to photosensitizer photophysics, and the implementation is available only upon request, which constrains reproducibility and extension.

By contrast, Dodds *et al.* integrated active learning with reinforcement learning (RL–AL) to enhance sample efficiency under multi-objective optimization oracles such as docking and ROCS, and released their implementation openly.³² While this framework effectively explores exploration–exploitation trade-offs, it is tailored to generic molecular design tasks and lacks photophysical objectives or physics-informed acquisition strategies specific to photosensitizer discovery.

Building upon these directions, our study establishes a unified active-learning (AL) framework that extends beyond task-specific implementations and explicitly targets photosensitizer discovery. The framework integrates three tightly coupled components: (i) a large-scale ML–xTB-calibrated dataset of 655 197 photosensitizer candidates labeled for T_1/S_1 , achieving sub-0.08 eV mean absolute error (MAE) and reducing computational cost by 99% compared with TD-DFT;³³ (ii) a hybrid acquisition strategy that combines ensemble-based

uncertainty estimation with a physics-informed objective function and an early-cycle diversity schedule, enabling balanced exploration and exploitation during iterative sampling; and (iii) a fully open-source, reproducible implementation that facilitates transparent benchmarking and community reuse.

Experimental benchmarks demonstrate that the proposed sequential AL strategy, which first explores chemical diversity before focusing on target regions, consistently outperforms static baselines by 15–20% in test-set MAE. By integrating data calibration, uncertainty-driven acquisition, and open implementation in a single workflow, this framework overcomes the methodological limitations of previous AL systems and provides a generalizable, data-efficient paradigm for photosensitizer discovery, offering immediate applications in solar fuels, photocatalysis, and optoelectronic materials design.

2 The unified active learning framework

2.1 Overview of the active learning workflow

Our active-learning platform for photosensitizer discovery consists of four main stages: (i) preparation of the molecular dataset and chemical space definition, (ii) training of graph neural network surrogates to predict key photophysical properties, (iii) molecule selection through complementary acquisition strategies, and (iv) validation *via* quantum-chemical calculations or experiments. This workflow enables iterative improvement of the surrogate model and efficient exploration of the vast molecular design space (Fig. 1). The following subsections describe the construction of the chemical space and the details of the active-learning protocol.

2.2 Design space generation

The construction of a chemically relevant and computationally tractable design space forms the foundation of our active learning framework. Traditional approaches relying solely on expert intuition or brute-force enumeration fail to address the dual challenges of chemical diversity and computational feasibility.^{34,35}

We combined Simplified Molecular-Input Line-Entry System (SMILES) data from numerous public molecular datasets to construct a unified library of 655 197 candidate photosensitizer molecules. Each source dataset was chosen because it contributes molecules with relevant excited-state or optical properties, thereby ensuring that our merged collection covers a broad range of photophysical characteristics. Starting from an initial seed set of 50 000 molecules, we expanded the library by integrating many diverse data sources (computational, experimental, and even patent-derived). We then predicted the lowest singlet and triplet energies (S_1 and T_1) for all candidates using our ML–xTB workflow, achieving DFT-level accuracy at 1% of the typical cost. We performed the analysis of the dataset in Fig. 3. (Full details of the datasets, including their names, references, and selection criteria, are provided in the SI (Text S5).)

2.2.1 ML–xTB pipeline. The ML–xTB workflow comprises three stages (Fig. 2):



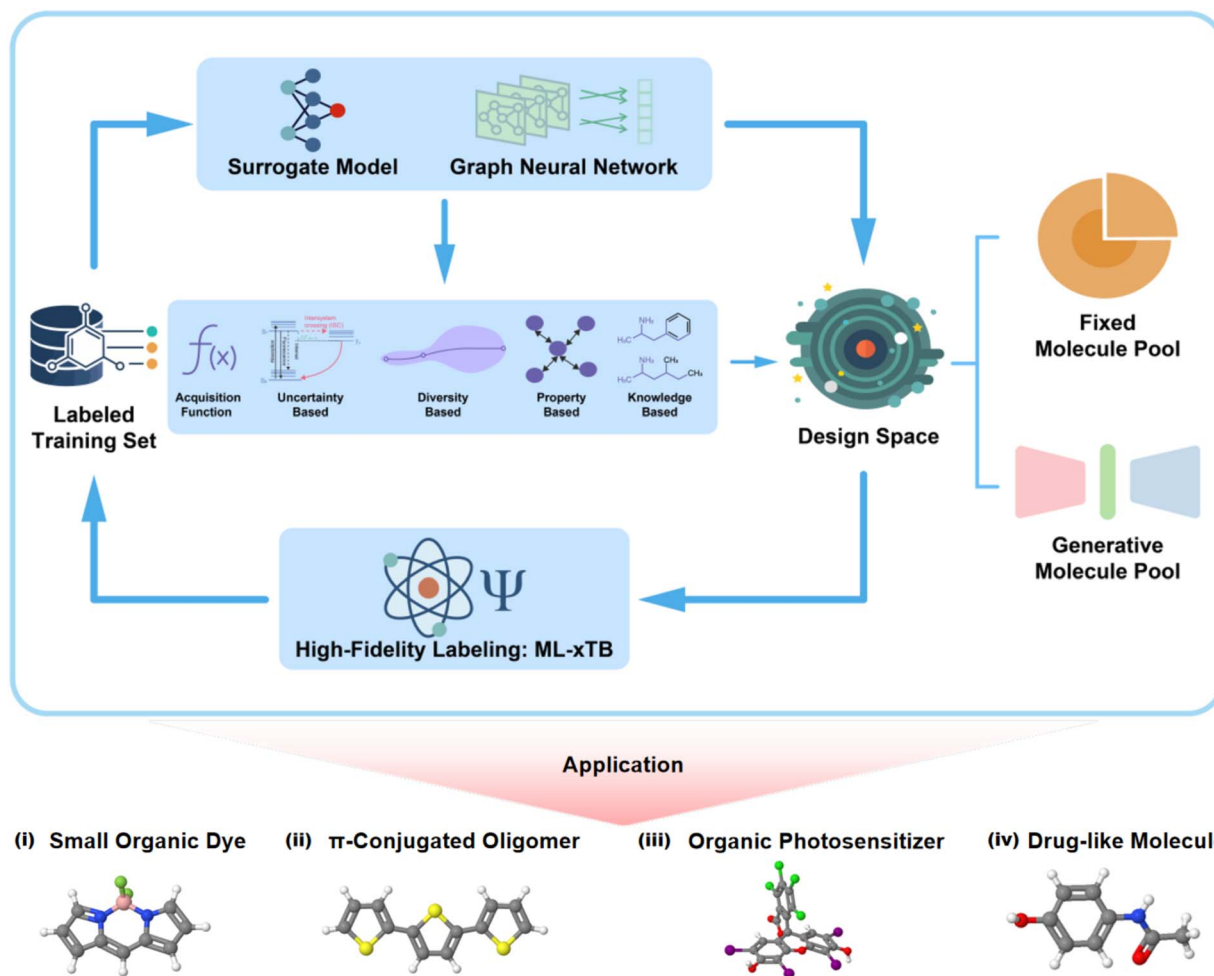


Fig. 1 Overview of the active learning framework for molecular design and property prediction. The workflow integrates a data-driven surrogate model with high-fidelity quantum-chemical labeling to accelerate molecular discovery. A graph neural network is trained on a labeled dataset to predict key electronic and photophysical properties. Candidate molecules are sampled from a combined design space that includes both a fixed molecule pool and a generative molecule pool, enabling exploration of known and novel chemical structures. Molecule selection for labeling is guided by four complementary strategies: uncertainty-based, diversity-based, property-based, and knowledge-based acquisition. High-accuracy electronic properties are then computed using ML-xTB, and the resulting data are iteratively added to the training set to refine the surrogate model. Representative applications include (i) small organic dyes, (ii) π -conjugated oligomers, (iii) organic photosensitizers, and (iv) drug-like molecules, illustrating the framework's generality across diverse molecular classes.

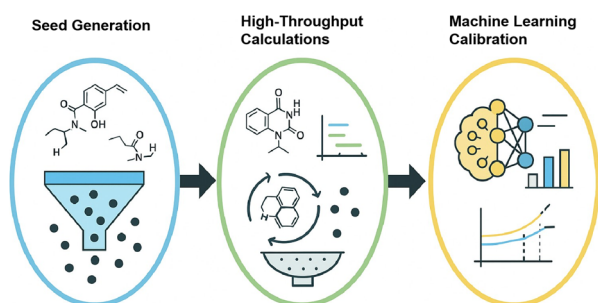


Fig. 2 ML-xTB pipeline for large-scale molecular property calculation. The process consists of initial seed generation from public databases and expert scaffolds, rapid pre-screening of excited-state energies using semi-empirical methods, and subsequent machine learning calibration to achieve DFT-level accuracy at substantially reduced computational cost.

1. Initial seed generation: A diverse set of 50 000 molecules was curated from public databases (PubChemQC,³⁶ QMspin³⁷) and expert-designed scaffolds (porphyrins, phthalocyanines). SMILES strings were standardized using RDKit, with stereochemistry and tautomer states normalized *via* Morgan fingerprint clustering (radius = 2, 1024 bits).

2. xTB-sTDA high-throughput calculations: Each molecule underwent geometry optimization and excited-state calculation using the geometry, frequency, noncovalent-tight binding (GFN2-xTB)³⁸ method combined with the simplified Tamm-Dancoff approximation (sTDA)³⁹ implemented in xtb (GFN2-xTB/xtb-sTDA):

$$S_1 = E_{\text{singlet}} - E_{\text{ground}} \quad (1)$$

$$T_1 = E_{\text{triplet}} - E_{\text{ground}} \quad (2)$$



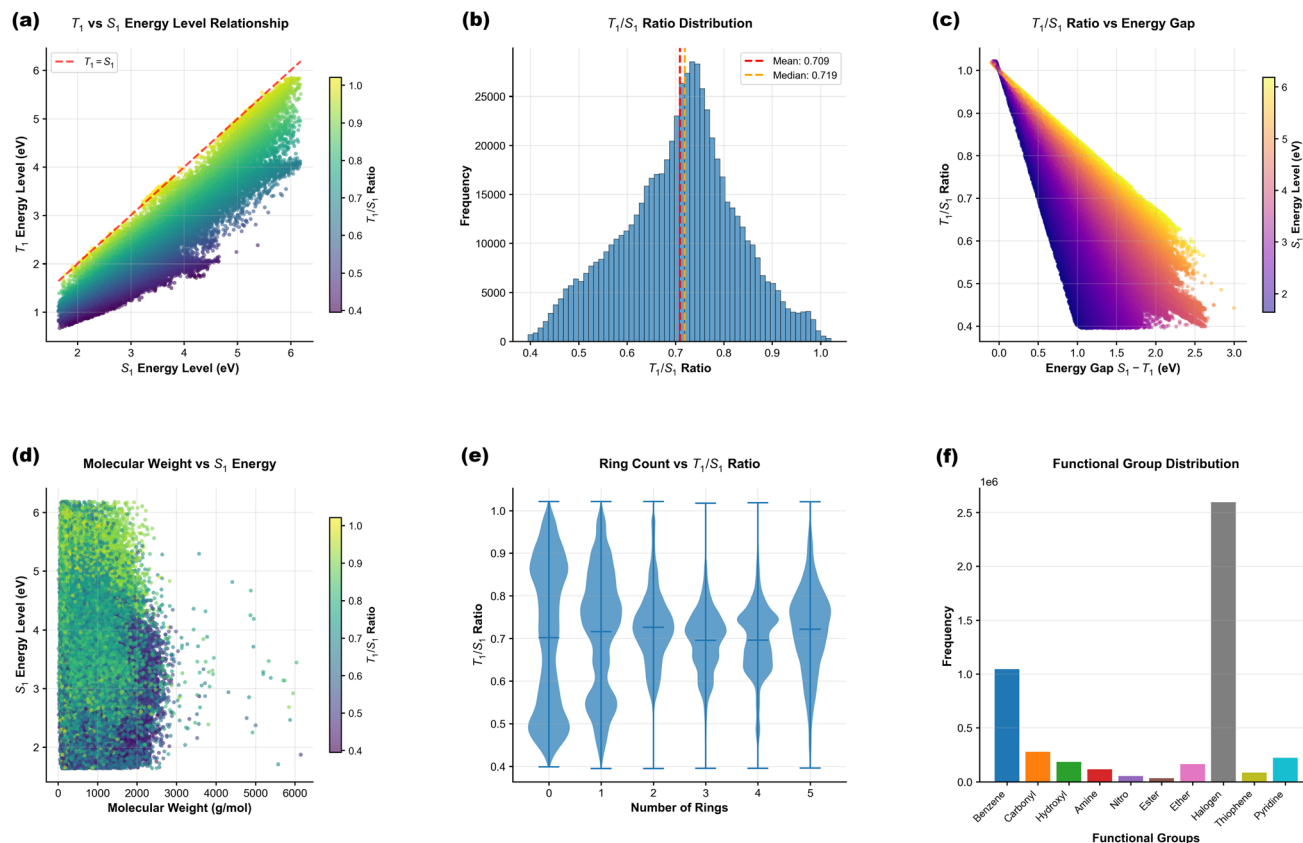


Fig. 3 Overview of key molecular and photophysical features in the unified active learning dataset. (a) T_1 vs. S_1 energy levels for 655 197 candidates show a near-linear correlation (colored by T_1/S_1 ratio), with most molecules below the $T_1 = S_1$ diagonal, highlighting typical singlet-triplet splitting and energetically favored structures. (b) Distribution of T_1/S_1 ratios, peaking near 0.7, provides a quantitative reference for selecting candidates with optimal photophysical properties. (c) T_1/S_1 ratio as a function of the $S_1 - T_1$ energy gap (colored by S_1), showing a triangular distribution: larger gaps yield lower T_1/S_1 ratios, revealing a fundamental structure-property trade-off. (d) S_1 energy as a function of molecular weight (colored by T_1/S_1 ratio), indicating broad chemical diversity and little direct dependence between S_1 and molecular weight. (e) Violin plots of T_1/S_1 ratio for different ring counts, showing robust distribution across core structures with only minor variation in highly fused systems. (f) Functional group statistics of the dataset: halogenated, aromatic, and carbonyl-containing molecules dominate, ensuring chemical diversity for generalizable model development.

$$\Delta E_{ST} = S_1 - T_1 \quad (3)$$

For the initial seed set (50 000 molecules), additional TD-DFT calculations (B3LYP/6-31+G(d), Gaussian 16) were performed on the xTB-optimized geometries to provide accurate reference values (details provided in SI).

3. Machine learning calibration: A 10-model ensemble of Chemprop Message Passing Neural Networks (Chemprop-MPNN) was trained to correct systematic errors between the 50 000 xTB-sTDA and TD-DFT calculations for the S_1 and T_1 excitations separately. Each network dynamically generated molecular representations from SMILES strings without static fingerprints, predicting state-specific errors:

$$\Delta E_S(i) = E_{S_1}^{\text{DFT}}(i) - E_{S_1}^{\text{xTB}}(i) \quad (4)$$

$$\Delta E_T(i) = E_{T_1}^{\text{DFT}}(i) - E_{T_1}^{\text{xTB}}(i) \quad (5)$$

The multitask loss function minimized during training was:

$$\mathcal{L} = \frac{1}{N} \sum_{i=1}^N [(\Delta E_S(i) - f_S(\mathbf{x}_i))^2 + (\Delta E_T(i) - f_T(\mathbf{x}_i))^2] \quad (6)$$

Where $f_S(\mathbf{x}_i)$ and $f_T(\mathbf{x}_i)$ are MPNN-predicted corrections for singlet and triplet excitations respectively.

The calibrated energies were then computed as:

$$E_{S_1}^{\text{corr}} = E_{S_1}^{\text{xTB}} + f_S(\mathbf{x}) \quad (7)$$

$$E_{T_1}^{\text{corr}} = E_{T_1}^{\text{xTB}} + f_T(\mathbf{x}) \quad (8)$$

for all molecules in the full dataset (655 197 molecules). We performed only xTB-sTDA calculations for the remaining molecules beyond the seed set and then applied the ML correction model. This calibration approach reduced the mean absolute error (MAE) from 0.23 eV (raw xTB) to 0.08 eV (ML-corrected) with respect to TD-DFT for the 50 000 molecules in the calibration set.

2.2.2 Dataset splitting and active learning protocol. An initial training set of 5000 molecules was randomly selected and kept consistent across all strategies. In each active learning



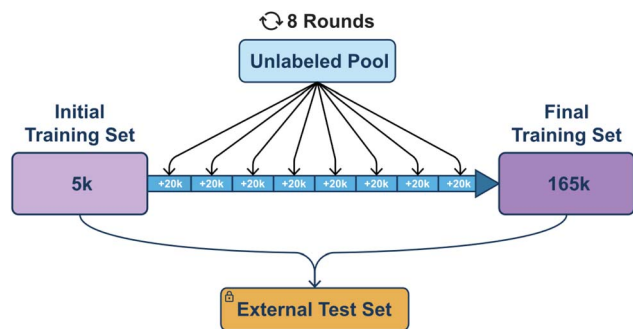


Fig. 4 Photosensitizer candidate data partition. A fixed external test set of molecules was first reserved from the complete dataset before any model training.

round, 20 000 additional molecules were sampled from the remaining pool, with a total of 8 rounds conducted for each acquisition strategy. This protocol enables model development and tuning on one portion of the data, while reserving a representative set of promising candidates for unbiased final evaluation (Fig. 4).

2.3 Surrogate model: Chemprop-MPNN

2.3.1 Rationale for model selection. The directed message-passing neural network (D-MPNN) from the Chemprop framework was selected as the surrogate model for its strong performance in molecular property prediction.^{40,41} (The same Chemprop architecture is also used in Section 2.1 as a Δ -learning calibration model that predicts the TD-DFT-xTB error; here, by contrast, it directly outputs absolute S_1 and T_1 energies as an surrogate model.)

This choice was driven by two key advantages: (1) Its explicit modeling of bond directionality (*e.g.*, single, double, conjugated) reduces noise from undirected representations, which is critical for capturing electronic transitions in photoactive molecules; and (2) its native ensemble support enables simultaneous quantification of model and data uncertainties, making it highly suitable for active learning strategies.

2.3.2 Uncertainty quantification. An ensemble of five independently trained D-MPNNs provided epistemic uncertainty estimates:

$$\text{Total variance} = \underbrace{\frac{1}{5} \sum_{k=1}^5 (\hat{y}_k - \bar{y})^2}_{\text{Model uncertainty}} \quad (9)$$

where $\bar{y}$ is the mean of the ensemble.

2.3.3 Bayesian hyperparameter optimization. Key hyperparameters were tuned *via* Gaussian process-based Bayesian optimization to minimize validation MAE:

$$\theta^* = \underset{\theta}{\operatorname{argmin}} \frac{1}{N_{\text{val}}} \sum_{i=1}^{N_{\text{val}}} \left(\left| S_1^{(i)} - \hat{S}_1^{(i)}(\theta) \right| + \left| T_1^{(i)} - \hat{T}_1^{(i)}(\theta) \right| \right) \quad (10)$$

During the active learning process, as we incrementally increased the training set size from 5k to 165k samples, we

conducted independent Bayesian hyperparameter optimizations at each data scale to determine the optimal model architectures. The hyperparameter search employed Root Mean Square Error (RMSE), the default evaluation metric in Chemprop, as the criterion for selecting the best configuration.

Our analysis showed that the optimal architecture did not converge to a single universal configuration; instead, the best-performing hyperparameters varied distinctly with the dataset size, including model depths (ranging from 3 to 6), hidden layer widths (hidden_size ranging from 1600 to 1900), and dropout rates (0.05–0.25). In medium-to-large data regimes (65–165k), shallow architectures (depth = 3), wider hidden layers (hidden_size = 1900), and lower dropout rates (0.05) consistently outperformed other configurations. Adopting these individually optimized hyperparameters at each training set size reduced the test RMSE by approximately 8–15% compared to an untuned baseline model (*e.g.*, depth = 3, hidden_size = 1200, dropout = 0.10). This stage-specific optimization strategy ensured optimal architecture selection at each phase of the active learning cycle, thereby eliminating potential bias from a fixed architecture and enabling fair comparisons across different acquisition strategies.

2.4 Acquisition strategies

Four strategies were systematically benchmarked to balance exploration and exploitation (Fig. 5):^{42,43}

2.4.1 Uncertainty sampling. For each molecule x , we compute the ensemble variance of the two targets— S_1 and T_1 —using eqn (9). The acquisition score is the sum of these two variances:

$$A_{\text{score}}(x) = \sigma_{S_1}^2(x) + \sigma_{T_1}^2(x), \quad (11)$$

where $\sigma_{S_1}^2(x)$ and $\sigma_{T_1}^2(x)$ are the predictive variances of S_1 and T_1 , respectively, obtained from the five-model ensemble. Candidates with the highest $A_{\text{score}}(x)$ probe regions where the surrogate is least confident, thereby accelerating error reduction.

2.4.2 Diversity-enhanced sampling. To avoid selecting many nearly identical molecules, we penalize intra-batch similarity:

$$A_{\text{score}}(\mathbf{x}) = A_{\text{base}}(x) \prod_{\mathbf{x}' \in B} \mathbb{I}(\text{Tanimoto}(\mathbf{x}, \mathbf{x}') < 0.6), \quad (12)$$

where B is the set of molecules already chosen for the current batch. The indicator function $\mathbb{I}(\cdot)$ discards a candidate if its radius-2, 2048 bit Morgan fingerprint yields a Tanimoto similarity greater than 0.6 to any member of B , thereby guaranteeing sufficient structural diversity while still prioritizing high-uncertainty points.

The Tanimoto similarity between two molecules $\mathbf{x}$ and $\mathbf{x}'$, based on their binary fingerprint vectors $f(\mathbf{x})$ and $f(\mathbf{x}')$, is defined as:

$$\text{Tanimoto}(\mathbf{x}, \mathbf{x}') = \frac{c}{a + b - c} \quad (13)$$



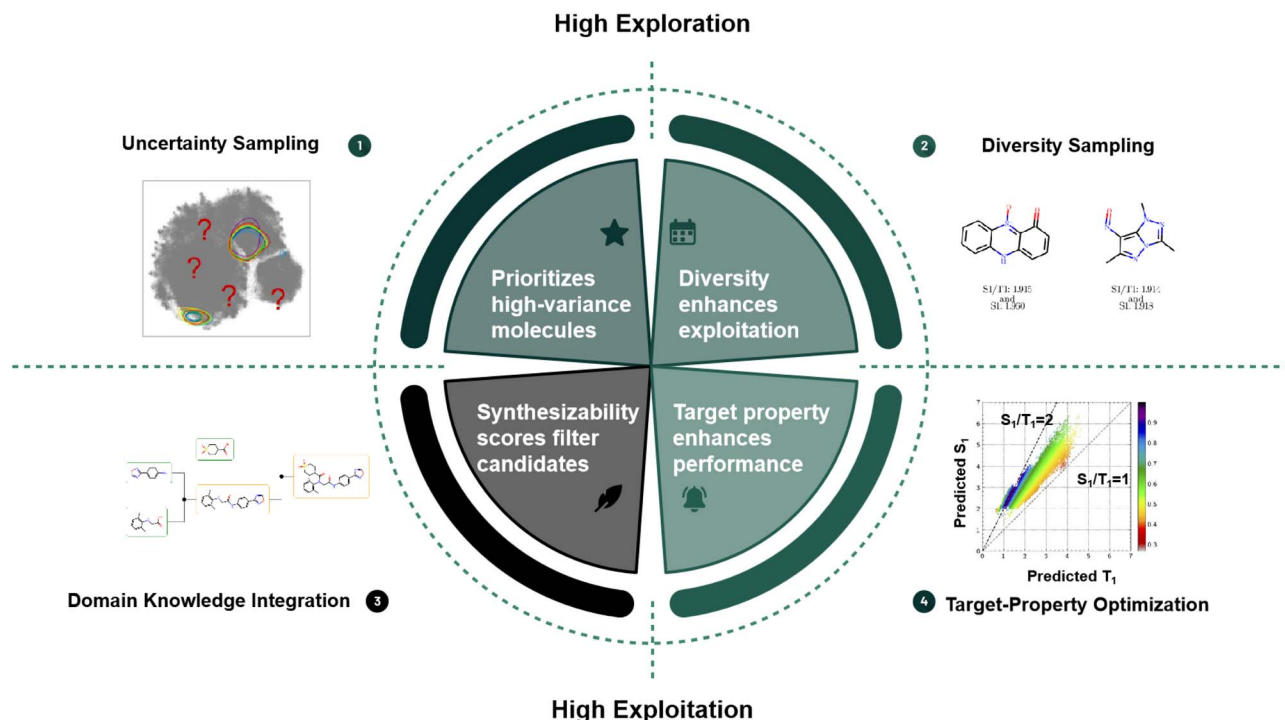


Fig. 5 Balancing exploration and exploitation in molecular selection. Four acquisition strategies are illustrated: uncertainty sampling prioritizes high-variance molecules; diversity sampling enhances exploration across chemical space; target-property optimization focuses on high-performance regions; and domain knowledge integration ensures synthetic feasibility. The combination of these strategies accelerates the identification of promising photosensitizer candidates.

where $a = \sum_i f_i(\mathbf{x})$ and $b = \sum_i f_i(\mathbf{x}')$ are the numbers of 'on' bits in each fingerprint, and $c = \sum_i f_i(\mathbf{x})f_i(\mathbf{x}')$ is the number of bits shared by both fingerprints. Here, $f(\mathbf{x})$ denotes the Morgan fingerprint of molecule $\mathbf{x}$; the Tanimoto similarity measures the overlap between the binary features of two molecules.

2.4.3 Target-property optimisation. We guide sampling toward molecules whose photophysical ratio

$$m(\mathbf{x}) = \frac{T_1(\mathbf{x})}{S_1(\mathbf{x})}$$

is close to either of two design targets: emitter and sensitizer.³³ The overall score is a weighted sum of the two targets, with tunable weights $w_e, w_s \geq 0$ satisfying $w_e + w_s = 1$ (default $w_e = w_s = 0.5$).

v1 — Probability-within-band kernel

$$A_{\text{score}}(\mathbf{x}) = w_e \int_{\tau_e - \epsilon}^{\tau_e + \epsilon} \mathcal{N}(m; \mu(\mathbf{x}), \sigma^2(\mathbf{x})) dm + w_s \int_{\tau_s - \epsilon}^{\tau_s + \epsilon} \mathcal{N}(m; \mu(\mathbf{x}), \sigma^2(\mathbf{x})) dm. \quad (14)$$

v2 — Exponential alignment kernel

$$A_{\text{score}}(\mathbf{x}) = w_e \exp[-\eta|m(\mathbf{x}) - \tau_e|] + w_s \exp[-\eta|m(\mathbf{x}) - \tau_s|]. \quad (15)$$

v3 — Expected-improvement kernel. An EI-style variant that accounts for correlation between T_1 and S_1 is derived in Text S3;

its numerical results are reported in Fig. S3 and therefore omitted here for brevity.

Detailed derivations, hyperparameter sweeps, and a head-to-head comparison of v1–v3 are provided in SI (Text S1–S3).

2.4.4 Symbol definitions (default values in parentheses).

- $m(\mathbf{x})$ — predicted ratio T_1/S_1 for molecule $\mathbf{x}$;
- $\mu(\mathbf{x}), \sigma^2(\mathbf{x})$ — mean and variance of the ensemble prediction for $m(\mathbf{x})$;
- $\tau_e = 0.5, \tau_s = 1.0$ — emitter-like and sensitizer-like targets;^{44,45}
- w_e, w_s — weights for the two targets (0.5, 0.5);
- ϵ — half-width of the tolerance band (0.05);
- η — selectivity parameter in v2 (5).

2.4.5 Domain knowledge integration. To ensure that candidates proposed by our active learning (AL) loop are experimentally viable, we introduced a synthetic-feasibility prior. The approach is inspired by the Retrosynthetic Accessibility Score (RAscore),⁴⁶ which predicts synthesizability from AiZynthFinder output.⁴⁷ Instead of using the original, domain-agnostic RAscore model, we trained a more expressive Chem-prop graph neural network on a photoswitch-specific dataset to obtain higher accuracy in the chemical space of interest.

2.7×10^5 emitter, sensitizer and analogue structures were subjected to AiZynthFinder. A molecule was labeled synthesized ($y = 1$) if the planner discovered at least one complete route to commercially available precursors within a user-defined search limit (maximum nstep retrosynthetic steps or t_{max} seconds). Molecules for which no route was found under these restrictions



were labeled non-synthesizable ($y = 0$). The resulting binary data set was used to train a Chemprop message-passing neural network that outputs a continuous probability $P_{\text{synth}}(\mathbf{x}) \in [0, 1]$. We refer to this domain-adapted score as the PhotoSynthScore.

During each AL iteration, we down-weight candidates that are unlikely to be synthesizable:

$$A_{\text{score}}(\mathbf{x}) = A_{\text{base}}(\mathbf{x}) \cdot \mathbb{I}(P_{\text{synth}}(\mathbf{x}) \geq 0.6), \quad (16)$$

where $A_{\text{base}}(\mathbf{x})$ is the score from Strategies 1–3 and $\mathbb{I}(\cdot)$ is the indicator function. Only molecules with $P_{\text{synth}} \geq 0.6$ pass the filter, focusing computational effort on candidates that are both high-performing and practically attainable. Empirically, incorporating the PhotoSynthScore improves search efficiency and maintains chemical realism throughout the discovery campaign.

3 Results and discussion

3.1 Sampling strategy benchmarking

Random sampling (baseline): To establish a fair reference, five independent random-sampling baselines were conducted under identical model architectures, hyperparameters, and labeling budgets. At each acquisition round, molecules were uniformly drawn from the unlabeled pool with the same batch size as in the active learning cycles. Averaged over five random seeds, the random baseline yielded a final external-test MAE of 0.083 eV. This result confirms that the observed performance improvement originates from the acquisition strategy besides data volume, demonstrating that the active learning framework effectively prioritizes more informative molecules for labeling.

Uncertainty sampling: Reduced MAE from 0.091 eV to 0.077 eV within 4 rounds, with performance plateauing at Round 8. Deepening hidden layers further improved accuracy across all test sets, particularly for large molecules (>20 non-H atoms).

Diversity-enhanced sampling: Introducing Tanimoto similarity thresholds (<0.6) increased heterocyclic system coverage by 25%, but aggressive thresholds (<0.4) raised RMSE by 15% due to oversampling of chemically irrelevant regions.

Target-property optimization: Focused screening of T_1/S_1 ratios (1 for sensitizers, 0.5 for emitters) reduced emitter MAE from 0.065 eV $\rightarrow$ 0.061 eV *versus* uncertainty sampling.

Synthetic feasibility integration: Threshold filtering ($P_{\text{synth}} \leq 0.6$) eliminated over 20% of all candidates, as these were judged to be impractical for synthesis.

3.2 Model performance and dataset size

As the size of the training dataset increases, the overall prediction accuracy of our model improves significantly, reflected by decreasing mean absolute error (MAE) and root mean square error (RMSE). However, this improvement demonstrates diminishing returns, indicating that once the dataset is sufficiently large, further additions yield limited performance gains. This trend is clearly depicted by the error curves across active learning rounds (Fig. 6), where initial rounds substantially reduce prediction errors, followed by a gradual plateau

indicating saturation of the model's capacity to extract useful information.

Most predictions achieve high accuracy within acceptable chemical precision thresholds. However, a small fraction of compounds exhibit notably larger errors, likely arising from unique or rare structural motifs not adequately represented in the training set. This highlights areas for future model refinement, particularly addressing the “long-tail” of challenging cases.

3.3 Impact of sampling strategies on data distribution and model generalization

We investigated how various active learning sampling strategies, uncertainty sampling, diversity sampling, target property optimization, and synthesizability-guided sampling, impact data set composition and model generalization.

Uncertainty sampling prioritizes molecules that the model finds the most ambiguous, thus rapidly improving the model precision by identifying informative samples. Diversity sampling aims to maximize coverage of chemical space, ensuring structural heterogeneity and preventing redundancy. Target-property optimization selects molecules with extreme property values, facilitating efficient exploration of high-performance regions, but potentially neglecting broader chemical diversity. Lastly, synthesizability-guided sampling incorporates synthetic feasibility, prioritizing practical compounds that are experimentally accessible.

These differences directly affect model training dynamics and generalization performance. Uncertainty sampling yields rapid initial reductions in prediction errors, demonstrating superior efficiency. Diversity sampling ensures broader chemical generalization but shows moderate accuracy improvements. Target-property sampling initially contributes less to global error reductions but significantly enhances predictions for extreme property cases. Synthesizability-guided strategies moderately slow accuracy gains due to conservative selections but ensure higher practical relevance.

We further summarized efficiency and error metrics for each strategy in Table 1, providing comprehensive comparisons on test MAE, and chemical space coverage. This analysis facilitates the selection of strategies that align with research priorities, balancing prediction accuracy, scope of the study, and practical applicability.

3.4 Portfolio strategies: multi-phase active learning and synthesizability constraints

Given the trade-offs above, an effective approach is to combine exploration and exploitation in a staged or simultaneous fashion—what we term a portfolio strategy. The idea is to first use uncertainty sampling to broadly train the model (exploration), and then gradually or abruptly shift toward selecting for the target property as the model becomes more reliable (exploitation). This latter approach dynamically balances exploration *vs.* exploitation within each iteration, rather than in separate phases.



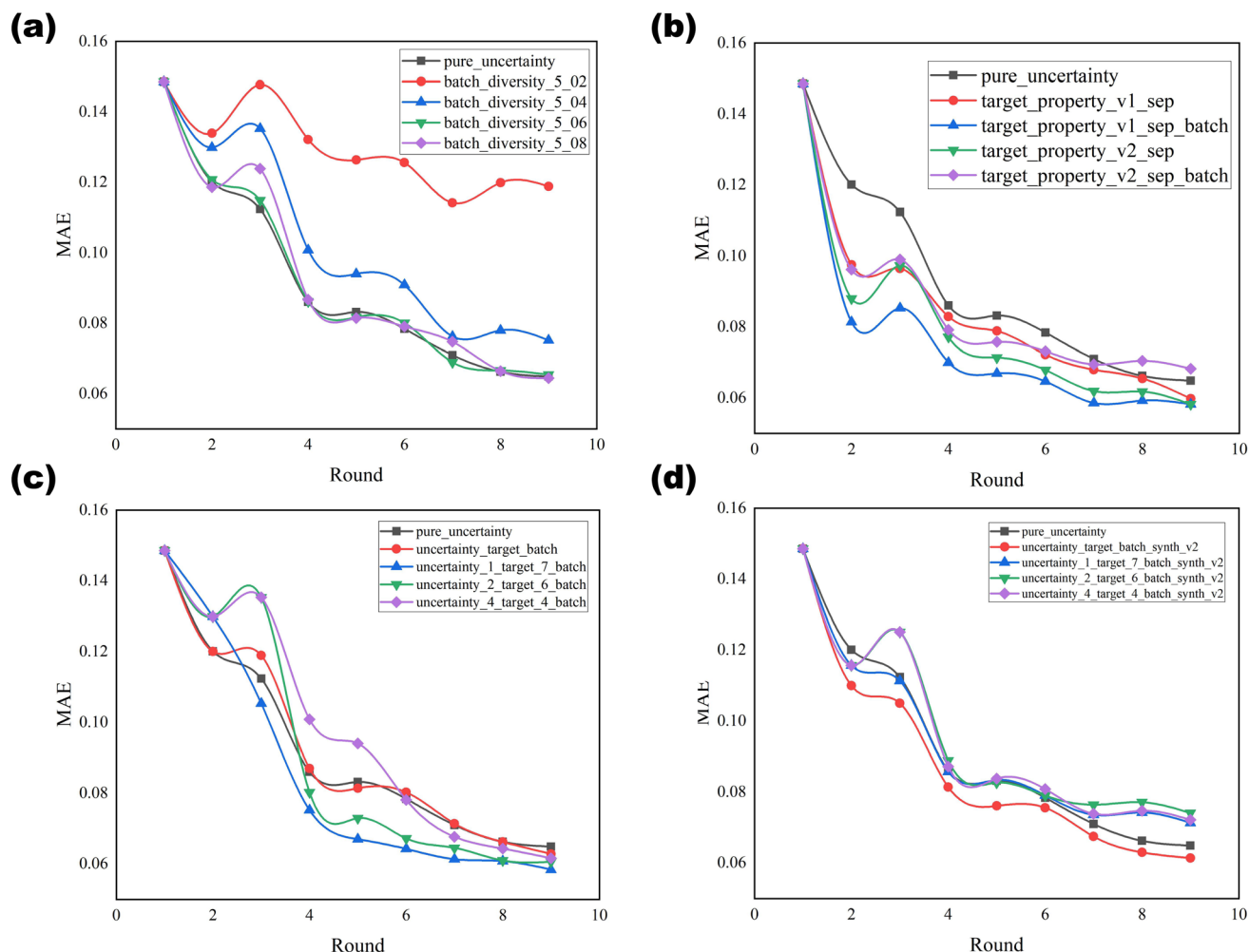


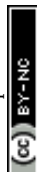
Fig. 6 Evolution of test-set MAE over active-learning round r for four representative acquisition strategies: (a) `uncertainty_batch_diversity_m_0τ` – Uncertainty-driven acquisition augmented by a diversity filter: a candidate is excluded if its Tanimoto similarity exceeds $\tau/10$ relative to at least m molecules already selected in the same batch; (b) `target_property_vφ_sep_batch` – Target-property-driven acquisition (version ϕ , see Section 2.4), combined with the diversity filter. Each batch split into half emitter-like and half sensitizer-like molecules; `target_property_vφ_sep` – Target-property-driven acquisition (version ϕ , see Section 2.4), without applying the diversity filter; (c) `uncertainty_a_target_b_batch` – Sequential portfolio strategy: the first a rounds employ uncertainty sampling, followed by b rounds of target-property sampling, both applying the diversity filter; `uncertainty_target_batch` – Weighted-sum portfolio strategy: uncertainty and target-property criteria are simultaneously considered in each round, applying the diversity filter; (d) `uncertainty_a_target_b_batch_synth_v2` – Same sequential portfolio strategy as (c), but incorporating an additional synthetic feasibility filter to exclude candidates predicted to be unsynthesizable; `uncertainty_target_batch_synth_v2` – Weighted-sum portfolio strategy (as above), augmented with the synthetic feasibility filter.

3.4.1 Effect of synthesizability filtering on portfolio performance. Before comparing the exploration–exploitation schedules, we consider the role of a synthesizability filter. This filter excludes candidates predicted to be synthetically

infeasible, adding a practical constraint to the selection process. Interestingly, we found that this constraint also has a pronounced effect on the learning dynamics. Both variants perform similarly in the early uncertainty-driven rounds (since

Table 1 Performance comparison of different active learning sampling strategies

Strategy	Initial MAE (eV)	Final MAE (eV)	Rounds to convergence	Chemical space coverage
Pure random (baseline)	0.149	0.083 ± 0.001	—	—
Pure uncertainty	0.149	0.077	8	Moderate
Uncertainty + batch diversity	0.149	0.078	9	High
Target property + batch diversity	0.149	0.062	9	Moderate (target-focused)
Portfolio strategy	0.149	0.059	9	High (balanced)
Portfolio + synthesizability	0.149	0.074	9	High (practically feasible)



those initial picks often tend to be relatively simple molecules or at least those the model is uncertain about, which may or may not be synthesizable).

Specifically, the filtered portfolio strategy shows a slight uptick or oscillation in MAE during the final rounds. After the switch to target-driven selection, the model begins choosing molecules that it predicts have extreme target values—some of these are likely unusual, highly complex structures that the model is less familiar with (and which might be chemically difficult to synthesize in reality). Incorporating such exotic compounds can momentarily worsen the model's overall performance: the newly added data might lie outside the domain where the model has predictive strength, leading to higher prediction errors for those points (and possibly disrupting the model's previously learned correlations). In other words, without the synthesizability filter, the model sometimes ventures into out-of-distribution regions in pursuit of high target property, and as a result the global MAE stops decreasing and even increases slightly in that phase. This phenomenon can be interpreted as a form of model extrapolation or overfitting issue—the model is essentially overextending into chemical space where its predictions are not reliable, analogous to how an overly aggressive exploitative strategy can mislead the model.

3.4.2 Balancing exploration and exploitation in multi-phase strategies. Finally, we assess the general effectiveness of the portfolio strategies and how different balances between exploration (uncertainty sampling) and exploitation (target-based selection) influence outcomes. The multi-phase approaches allow us to tune the exploration–exploitation trade-off by deciding how many rounds to devote to each.

One notable advantage of the hybrid strategy is that it avoids an abrupt transition that might destabilize the model. In the sequential strategies, we sometimes see a kink or change in the MAE trend at the point of switching from uncertainty to target mode—if the switch happens too early, the model could struggle (as discussed, a slight MAE rise can occur if the model isn't ready to accurately handle the exploitation picks). The continuous strategy softens this by always maintaining a mix; effectively, it performs a dynamic rebalancing: as the model's confidence grows, more of the selected batch inherently contributes to exploitation (since fewer points will be at high uncertainty, the focus naturally shifts to high predicted property, without ever entirely ignoring uncertainty). This dynamic adjustability is a strong point of the hybrid method. The only slight drawback observed is that the hybrid strategy can be a bit slower to find the very top-performing molecules compared to a full-on exploitation in later rounds. Since it's never selecting only target-optimal candidates, it might miss a few opportunities to immediately test the absolute top predicted molecule in favor of an uncertain one. However, in practice this seems to be a minor penalty—the hybrid still discovers high-performing molecules throughout, just interspersed with exploratory picks. In exchange, it maintains the lowest MAE curve among the methods, indicating it never compromises the model's learning too much in pursuit of the objective.

4 Conclusions

In this work, we present a unified computational framework to tackle the main challenges in photosensitizer discovery by introducing three key innovations:

- We developed a hybrid pipeline that combines fast semi-empirical quantum calculations with machine learning correction, enabling us to generate a large dataset of 655 197 diverse candidate molecules. This approach keeps the speed of computation high while ensuring the accuracy of quantum chemistry.
- We designed a sequential active learning strategy that separates the exploration and exploitation phases. First, the model explores chemical space broadly using uncertainty-driven sampling to find new promising areas, and then focuses on finding molecules with the desired T_1/S_1 energy ratios (1 for sensitizers, 0.5 for emitters).

Our open dataset and modular workflow offer a useful resource for the research community to speed up the development of new energy materials. Through detailed analysis, we show that an active learning approach which combines uncertainty sampling, target-driven optimization, diversity filters, and synthesizability constraints is most effective. This balanced strategy helps us find high-performing molecules faster, while also making sure the predictive model works well on different types of molecules. Adding diversity filters prevents the model from overfitting to very similar compounds, and considering synthesizability ensures the candidates found are likely to be made in real experiments.

While our model works well for small and moderate-sized molecules,⁴⁸ it is less accurate for larger molecules with highly delocalized electronic structures or strong long-range interactions. This is mainly because the current graph neural networks are better at capturing local structure, and we only use 2D information rather than 3D shapes. For molecules whose properties depend on their 3D arrangement, this can limit our model's performance.

To overcome these issues, future work could explore hybrid architectures that insert multi-head transformer self-attention directly into the message-passing layers of a GNN, going beyond the readout functions currently used in Chemprop, thus capturing both local chemical environments and long-range interatomic dependencies. Adding 3D geometric features, such as distances and angles between atoms, will also help the model better predict properties that depend on molecular shape.⁴⁹

We also see value in combining multiple types of molecular information, such as electron density and orbital properties, and predicting a range of photophysical properties at once. This multi-task and multimodal approach can make the model stronger and easier to interpret.

Additionally, looking forward, computational methods for molecular discovery have evolved significantly, moving from simple virtual screenings to fully automated experimental platforms. Initially, high-throughput virtual screening (HTVS) allowed researchers to computationally evaluate thousands of



candidate molecules from static libraries to quickly identify promising leads.⁹ Next, active learning methods greatly improved efficiency by selectively and iteratively choosing the most informative candidates, thereby reducing computational cost and speeding up discovery. More recently, a new generation of closed-loop discovery systems has emerged, combining active learning with fully automated laboratories that autonomously synthesize, test, and validate molecular candidates in continuous cycles.⁵⁰ By incorporating these successive advancements into photosensitizer discovery, future workflows could efficiently and automatically discover effective molecules for practical applications in energy and medicine.

In parallel, a promising direction for future work is to extend our framework from selecting candidates in fixed libraries to generating new molecular structures from scratch. This can be achieved by combining the trained surrogate model with generative techniques such as variational autoencoders, graph-based molecular generators, language-model-driven design, or fragment-based human-in-the-loop strategies. These methods would allow the model to propose novel compounds tailored to target singlet–triplet energy ratios and other design constraints, while exploring chemical space beyond the initial dataset.

For active learning, specifically, adaptive strategies, those that dynamically adjust sampling rules throughout different learning stages, can further improve efficiency and generalization. By smoothly moving between uncertainty, diversity, and target-based sampling, the model can learn more with fewer experiments.

In summary, by applying smart active learning, we were able to further increase the accuracy and usefulness of our method. These steps helped bridge the gap between fast ML predictions and high-accuracy quantum methods, making our framework a reliable tool for discovering new photoactive molecules. Our work highlights the importance of balancing exploration and exploitation, including chemical diversity and synthesizability, and adjusting learning strategies as the model grows, laying a practical foundation for future AI-driven molecular discovery.

Author contributions

Yizhe Chen conceptualized the study, developed the overall framework, implemented the active learning algorithms, and drafted the manuscript. Shomik Verma provided the datasets, performed the implementation of the computational method. Kevin P. Greenman contributed to the algorithm design within Chemprop and provided critical analysis and suggestions regarding the active learning methodologies and reviewed and edited the manuscript. Haoyu Yin, Zhihao Wang, and Lanjing Wang assisted with data visualization, figure preparation, and content curation. Jiali Li, Rafael Gómez-Bombarelli, Aron Walsh, and Xiaonan Wang supervised the project, reviewed and edited the manuscript, and provided oversight throughout the study.

Conflicts of interest

The authors declare no competing interests.

Data availability

The curated dataset of 655 197 photosensitizer candidates and the complete active-learning code are available at <https://github.com/jiali1025/>

A General Active learning framework for MoleDesign. The dataset is released under a CC-BY-4.0 licence and the code under an MIT licence. No further restrictions apply.

Supplementary information (SI): additional methodological details, acquisition-function formulations, computational protocols, and a description of the unified dataset supporting this work. See DOI: <https://doi.org/10.1039/d5sc05749c>.

Acknowledgements

This work is supported by the National Key R&D Program of China (No. 2022ZD0117501), the Scientific Research Innovation Capability Support Project for Young Faculty (ZYGXQJNSKYCXNLZCXM-E7) and Tsinghua University Initiative Scientific Research Program. K. P. G. was supported by the National Science Foundation Graduate Research Fellowship Program under Grant No. 1745302 and the DARPA Accelerated Molecular Discovery (AMD) program under contract HR00111920025. The authors acknowledge additional support from collaborating institutions and funding agencies.

Notes and references

- 1 K. T. Butler, D. W. Davies, H. Cartwright, O. Isayev and A. Walsh, *Nature*, 2018, **559**, 547–555.
- 2 Y. Cai, T. Chai, W. Nguyen, J. Liu, E. Xiao, X. Ran, Y. Ran, D. Du, W. Chen and X. Chen, *Signal Transduction Targeted Ther.*, 2025, **10**, 115.
- 3 T. Froitzheim, S. Grimme and J.-M. Mewes, *J. Chem. Theory Comput.*, 2022, **18**, 7702–7713.
- 4 H. Kim, Y. R. Lee, H. Jeong, J. Lee, X. Wu, H. Li and J. Yoon, *Smart Mol.*, 2023, **1**, e20220010.
- 5 F. Hu, S. Xu and B. Liu, *Adv. Mater.*, 2018, **30**, 1801350.
- 6 J. P. Janet, S. Ramesh, C. Duan and H. J. Kulik, *ACS Cent. Sci.*, 2020, **6**, 513–524.
- 7 S. Xu, J. Li, P. Cai, X. Liu, B. Liu and X. Wang, *J. Am. Chem. Soc.*, 2021, **143**, 19769–19777.
- 8 K. Chen, X. Zhang, J. Wang, D. Li, T. Hou, W. Yang and Y. Kang, *Chem. Sci.*, 2025, **16**, 14698–14709.
- 9 R. Gómez-Bombarelli, J. Aguilera-Iparraguirre, T. D. Hirzel, D. Duvenaud, D. Maclaurin, M. A. Blood-Forsythe, H. S. Chae, M. Einzinger, D.-G. Ha, T. Wu, *et al.*, *Nat. Mater.*, 2016, **15**, 1120–1127.
- 10 P. Xu, X. Ji, M. Li and W. Lu, *npj Comput. Mater.*, 2023, **9**, 42.
- 11 X. Kang, Z. Du, S. Yang, M. Liang, Q. Liu and J. Qi, *Smart Mol.*, 2024, **2**, e20240033.
- 12 K. M. Jablonka, G. M. Jothiappan, S. Wang, B. Smit and B. Yoo, *Nat. Commun.*, 2021, **12**, 2312.
- 13 M. Sumita, X. Yang, S. Ishihara, R. Tamura and K. Tsuda, *ACS Cent. Sci.*, 2018, **4**, 1126–1133.
- 14 X. Li, P. M. Maffettone, Y. Che, T. Liu, L. Chen and A. I. Cooper, *Chem. Sci.*, 2021, **12**, 10742–10754.



- 15 J. Moon, W. Beker, M. Siek, J. Kim, H. S. Lee, T. Hyeon and B. A. Grzybowski, *Nat. Mater.*, 2024, **23**, 108–115.
- 16 Y. Zhao, Q. Liu, J. Du, Q. Meng and L. Zhang, *Smart Mol.*, 2023, **1**, e20230012.
- 17 A. Merchant, S. Batzner, S. S. Schoenholz, M. Aykol, G. Cheon and E. D. Cubuk, *Nature*, 2023, **624**, 80–85.
- 18 B. Settles, *Active Learning Literature Survey* publisher (University of Wisconsin–Madison, Computer Sciences Technical Report 1648, 2009.
- 19 H. Chun, J. R. Lunger, J. K. Kang, R. Gómez-Bombarelli and B. Han, *npj Comput. Mater.*, 2024, **10**, 246.
- 20 C. Duan, A. Nandy, G. G. Terrones, D. W. Kastner and H. J. Kulik, *JACS Au*, 2022, **3**, 391–401.
- 21 R. Ding, J. Liu, K. Hua, X. Wang, X. Zhang, M. Shao, Y. Chen and J. Chen, *Sci. Adv.*, 2025, **11**, eadr9038.
- 22 M. Kim, Y. Kim, M. Y. Ha, E. Shin, S. J. Kwak, M. Park, I.-D. Kim, W.-B. Jung, W. B. Lee, Y. Kim, *et al.*, *Adv. Mater.*, 2023, **35**, 2211497.
- 23 L. Wang, Z. Zhou, X. Yang, S. Shi, X. Zeng and D. Cao, *Drug Discovery Today*, 2024, 103985.
- 24 T. Yin, G. Panapitiya, E. D. Coda and E. G. Saldanha, *J. Cheminf.*, 2023, **15**, 105.
- 25 H. H. Loeffler, S. Wan, M. Klahn, A. P. Bhati and P. V. Coveney, *J. Chem. Theory Comput.*, 2024, **20**, 8308–8328.
- 26 B. Cree, M. K. Bieniek, S. Amin, A. Kawamura and D. J. Cole, *Digital Discovery*, 2025, **4**, 438–450.
- 27 P. Shetty, A. Adeboye, S. Gupta, C. Zhang and R. Ramprasad, *Chem. Mater.*, 2024, **36**, 7676–7689.
- 28 S. Thaler, F. Mayr, S. Thomas, A. Gagliardi and J. Zavadlav, *npj Comput. Mater.*, 2024, **10**, 86.
- 29 L. Kavalsky, V. I. Hegde, B. Meredig and V. Viswanathan, *Digital Discovery*, 2024, **3**, 999–1010.
- 30 D. Buterez, J. P. Janet, S. J. Kiddle, D. Oglic and P. Lió, *Nat. Commun.*, 2024, **15**, 1517.
- 31 P. Yoo, D. Bhowmik, K. Mehta, *et al.*, *Sci. Rep.*, 2023, **13**, 20031.
- 32 M. Dodds, J. Guo, T. Löhr, A. Tibo, O. Engkvist and J. P. Janet, *Chem. Sci.*, 2024, **15**, 4146–4160.
- 33 S. Verma, M. Rivera, D. O. Scanlon and A. Walsh, *J. Chem. Phys.*, 2022, **156**(13), 134116.
- 34 G. Schneider and U. Fechner, *Nat. Rev. Drug Discovery*, 2005, **4**, 649–663.
- 35 L. Ruddigkeit, R. Van Deursen, L. C. Blum and J.-L. Reymond, *J. Chem. Inf. Model.*, 2012, **52**, 2864–2875.
- 36 M. Nakata, T. Shimazaki and H. Nakai, *J. Chem. Inf. Model.*, 2017, **57**, 1300–1308.
- 37 M. Schwilk, D. N. Tahchieva and O. A. von Lilienfeld, The QMspin data set: Several thousand carbene singlet and triplet state structures and vertical spin gaps computed at MRCISD+Q-F12/cc-pVDZ-F12 level of theory, *Materials Cloud Archive*, 2020, DOI: [10.24435/materialscloud:2020.0051/v1](https://doi.org/10.24435/materialscloud:2020.0051/v1).
- 38 C. Bannwarth, S. Ehlert and S. Grimme, *J. Chem. Theory Comput.*, 2019, **15**, 1652–1671.
- 39 S. Grimme and C. Bannwarth, *J. Chem. Phys.*, 2016, **145**(5), 054143.
- 40 E. Heid, K. P. Greenman, Y. Chung, S.-C. Li, D. E. Graff, F. H. Vermeire, H. Wu, W. H. Green and C. J. McGill, *J. Chem. Inf. Model.*, 2024, **64**, 9–17.
- 41 J. Westermayr and P. Marquetand, *Chem. Rev.*, 2020, **121**, 9873–9926.
- 42 M. Martyka, L. Zhang, F. Ge, Y.-F. Hou, J. Jankowska, M. Barbatti and P. O. Dral, *npj Comput. Mater.*, 2025, **11**, 1–12.
- 43 A. Nigam, R. Pollice, G. Tom, K. Jorner, J. Willes, L. Thiede, A. Kundaje and A. Aspuru-Guzik, *Adv Neural Inf Process Syst*, 2023, **36**, 3263–3306.
- 44 L. Naimovičius, P. Bharmoria and K. Moth-Poulsen, *Mater. Chem. Front.*, 2023, **7**, 2297–2315.
- 45 J. L. Weber, E. M. Churchill, S. Jockusch, E. J. Arthur, A. B. Pun, S. Zhang, R. A. Friesner, L. M. Campos, D. R. Reichman and J. Shee, *Chem. Sci.*, 2021, **12**, 1068–1079.
- 46 A. Thakkar, V. Chadimová, E. J. Bjerrum, O. Engkvist and J.-L. Reymond, *Chem. Sci.*, 2021, **12**, 3339–3349.
- 47 S. Genheden, A. Thakkar, V. Chadimová, J.-L. Reymond, O. Engkvist and E. Bjerrum, *J. Cheminf.*, 2020, **12**, 70.
- 48 R. Gómez-Bombarelli, J. N. Wei, D. Duvenaud, J. M. Hernández-Lobato, B. Sánchez-Lengeling, D. Sheberla, J. Aguilera-Iparraguirre, T. D. Hirzel, R. P. Adams and A. Aspuru-Guzik, *ACS Cent. Sci.*, 2018, **4**, 268–276.
- 49 X. Fang, L. Liu, J. Lei, D. He, S. Zhang, J. Zhou, F. Wang, H. Wu and H. Wang, *Nat. Mach. Intell.*, 2022, **4**, 127–134.
- 50 T. Wu, S. Kheiri, R. J. Hickman, H. Tao, T. C. Wu, Z.-B. Yang, X. Ge, W. Zhang, M. Abolhasani, K. Liu, *et al.*, *Nat. Commun.*, 2025, **16**, 1473.

