

# LEARNING OVER MOLECULAR CONFORMER ENSEMBLES: DATASETS AND BENCHMARKS

Yanqiao Zhu<sup>✉</sup> Jeehyun Hwang<sup>✉</sup> Keir Adams<sup>✉</sup> Zhen Liu<sup>✉</sup> Bozhao Nan<sup>✉</sup>  
 Brock Anton Stenfors<sup>✉</sup> Yuanqi Du<sup>✉</sup> Jatin Chauhan<sup>✉</sup> Olaf Wiest<sup>✉</sup>  
 Olexandr Isayev<sup>✉</sup> Connor W. Coley<sup>✉</sup> Yizhou Sun<sup>✉</sup> Wei Wang<sup>✉</sup>  
<sup>✉</sup>UCLA <sup>✉</sup>MIT <sup>✉</sup>CMU <sup>✉</sup>Notre Dame <sup>✉</sup>Cornell

✉ Primary contact: yzhu@cs.ucla.edu

🏠 Project homepage: <https://github.com/SXKDZ/MARCEL>

## ABSTRACT

Molecular Representation Learning (MRL) has proven impactful in numerous biochemical applications such as drug discovery and enzyme design. While Graph Neural Networks (GNNs) are effective at learning molecular representations from a 2D molecular graph or a single 3D structure, existing works often overlook the flexible nature of molecules, which continuously interconvert across conformations via chemical bond rotations and minor vibrational perturbations. To better account for molecular flexibility, some recent works formulate MRL as an ensemble learning problem, focusing on explicitly learning from a set of conformer structures. However, most of these studies have limited datasets, tasks, and models. In this work, we introduce the first Molecular AR Conformer Ensemble Learning (MARCEL) benchmark to thoroughly evaluate the potential of learning on conformer ensembles and suggest promising research directions. MARCEL includes four datasets covering diverse molecule- and reaction-level properties of chemically diverse molecules including organocatalysts and transition-metal catalysts, extending beyond the scope of common GNN benchmarks that are confined to drug-like molecules. In addition, we conduct a comprehensive empirical study, which benchmarks representative 1D, 2D, and 3D MRL models, along with two strategies that explicitly incorporate conformer ensembles into 3D models. Our findings reveal that direct learning from an accessible conformer space can improve performance on a variety of tasks and models.

## 1 INTRODUCTION

Recent years have seen the emergence of Molecular Representation Learning (MRL) as a promising approach for modeling molecules with machine learning. In the typical formulation, MRL maps discrete molecular objects to continuous features in a data-driven manner, encoding complex chemical structures into representation vectors that can subsequently be utilized in different downstream tasks. In particular, MRL now underpins a variety of biochemical applications spanning molecular property prediction to the design of novel drug candidates [1–3].

Traditional approaches often encode chemical compounds with fingerprints, such as extended-connectivity fingerprints [4, 5], which indicate the existence of certain substructures as binary bits in a fixed-length sequence. Such line-based representations are concise and efficient, but have limited expressive power and have difficulty in capturing 3D structural information such as bonding geometries and global shapes, which can be important for analyzing molecular properties and chemical reactivity [6, 7]. Recently, Graph Neural Networks (GNNs) have become an increasingly popular method of learning molecular representations by treating molecules as graph-structured objects. Existing GNN models for MRL can be broadly classified into two categories: 2D topological models [8–11] and 3D geometric models [12–17]. 2D GNNs typically model the molecular connectivity as a flat 2D graph with atoms as nodes and bonds as edges, learning representations of chemical environments by iteratively passing messages between neighboring atoms. Although powerful in the absence of structural information, 2D GNNs may fail to capture key conformational effects or stereochemical properties like chirality [18, 19], which is critical for modeling molecular interactions in areas such as

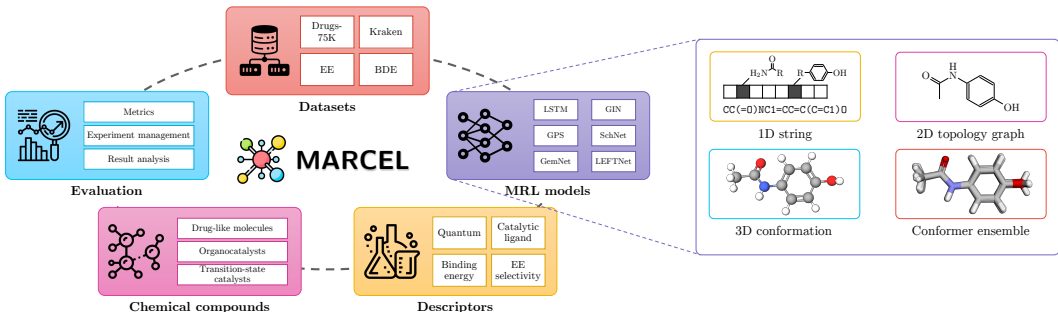


Figure 1: We present a MARCEL benchmark that comprehensively evaluates the potential of learning on conformer ensembles across a diverse set of molecules, datasets, and models.

drug design or chemical catalysis. Conversely, 3D GNNs are designed to model molecular conformers (conformations), which describe the structure of molecules in 3D space. Thus, these models have found widespread adoption for modeling electronic properties, predicting conformer energies and forces, and scoring interactions between ligands and proteins, amongst other applications.

In almost all applications, benchmarks, and demonstrations, 3D GNN models focus on encoding *individual* conformer structures. It is critical to recognize that in reality molecules are not rigid, static objects; rather, thermodynamically-permissible rotations of chemical bonds, small vibrational motions, and dynamic intermolecular interactions cause molecules to continuously convert between different conformations [20]. As a consequence, many experimentally observable chemical properties depend on the full distribution of thermodynamically-accessible conformers. For example, a molecule needs to be arranged into a particular pose to bind to a target protein, and this binding conformation changes depending on the dynamic interaction between the molecule and the target [21]. Also, it is often challenging to determine *a priori* the conformers that predominantly contribute to molecular properties without doing prohibitively expensive simulations. Therefore, a natural question arises: can we leverage the *collective* power of many different conformer structures lying on the local minima of the potential energy surface, also known as the *conformer ensemble*, to improve MRL models?

As shown by the empirical evidence from various studies, learning from an explicit conformer ensemble can prove to be advantageous for many tasks, including property and energy prediction [22–24], key conformer pose identification [25], and RNA sequence design [26]. However, these studies have been mostly confined to small-scale datasets, a limited set of tasks, and a restricted set of model architectures. As a result, it remains unclear (1) to what extent 2D GNNs can implicitly model molecular flexibility and (2) whether the *explicit* encoding of conformer ensembles can improve the performance of 3D models that traditionally encode only one single conformer.

In this paper, we present the first Molecular AConformer Ensemble Learning (MARCEL) benchmark. It covers a diverse range of chemical space (Figure 1), which focuses on four chemically-relevant tasks for both molecules and reactions, with an emphasis on Boltzmann-averaged properties of conformer ensembles computed at the Density-Functional Theory (DFT) level. Our datasets encompass a variety of compounds with high-quality conformers, including organocatalysts and transition-metal catalysts, extending beyond the scope of conventional GNN benchmarks which are often restricted to drug-like molecules. Moreover, we implement a benchmark suite that enables extensive empirical studies across representative 1D, 2D, and 3D models. We further explore the advantages of leveraging conformer ensembles through two straightforward strategies: (1) augmenting training samples by randomly selecting one conformer from the ensemble for each molecule and (2) applying an explicit multi-instance ensemble learning layer, which aggregates individual conformer embeddings.

Our experimental results confirm the potential effectiveness of incorporating conformer ensembles in MRL, highlighting the improvements over conventional single-conformation 3D networks. However, it is important to understand the heterogeneity of outcomes based on different dataset characteristics, task objectives, and model choices. Our investigation yields three key findings: (1) Leveraging molecular conformers by incorporating explicit set encoders, as a part of conformer ensemble learning strategies, can improve single-conformer 3D MRL models performance. (2) Data augmentation through conformer sampling may offer potential benefits, evidenced by improved results in the BDE dataset, suggesting a method to increase model robustness against imprecise structures. (3) Model selection for MRL depends on dataset sizes and tasks, with traditional 1D fingerprints and 2D models preferred for smaller datasets and 3D models for larger or reaction-focused tasks.

## 2 PROBLEM FORMULATION

We represent a 2D molecular graph as a tuple  $G = (\mathcal{V}, \mathcal{E}, \mathbf{X}, \mathbf{W})$ , where  $\mathcal{V} = \{v_i\}_{i=1}^{|\mathcal{V}|}$  is the node set with each node corresponding to an atom, and  $\mathcal{E} \subseteq \mathcal{V} \times \mathcal{V}$  is the edge set representing chemical bonds as edges between nodes. Further,  $\mathbf{X} \in \mathbb{R}^{d_v \times |\mathcal{V}|}$  contains vector attributes for each node, and  $\mathbf{W} \in \mathbb{R}^{d_w \times |\mathcal{E}|}$  contains attributes for each edge. When modeling chemical reactions, we represent a molecule-molecule complex as a pair of graphs  $(G_1, G_2)$ . In this case, the conformation describes the combined structure of the interacting molecules. For a given molecule or molecular complex, we assume that its geometry can be effectively characterized by a representative set of discrete, sampled conformers from the thermodynamically-accessible conformer distribution. Formally, this set can be denoted as  $\mathcal{C} = \{C_i\}_{i=1}^{|\mathcal{C}|}$ , where  $C_i \in \mathbb{R}^{|\mathcal{V}| \times 3}$  represents one conformer structure in 3D space. In reality, the conformer distribution is continuous;  $\mathcal{C}$  in our study contains representative samples of the infinite set. Each conformer in the sampled ensemble is associated with a statistical weight given by

$$p_i = \frac{\exp\left(-\frac{e_i}{k_B T}\right)}{\sum_j \exp\left(-\frac{e_j}{k_B T}\right)},$$

which corresponds to its probability under experimental conditions. Here,  $e_i$  is the energy of the conformer  $C_i$ ,  $k_B$  is the Boltzmann constant, and  $T$  is the temperature. Notably,  $p_i$  is not prior information to the models analyzed in this benchmark. Rather, we use a discrete approximation of  $p_i$  to compute the ground-truth labels for our regression tasks.

## 3 DATASETS AND TASKS

**MARCEL** contains four small-to-large-scale datasets involving nine regression tasks with considerably diverse chemistry. **Drugs-75K** and **Kraken** focus on molecular properties, while **EE** and **BDE** focus on reaction-centric properties. **MARCEL** includes molecules with high structural flexibility, evidenced by an average number of rotatable bonds exceeding 5. Table 1 summarizes the datasets.

**Drugs-75K** is a subset of the **GEOM-Drugs** [27] dataset, which includes 75,099 molecules with at least 5 rotatable bonds. For each molecule, **Auto3D** [28] is used to generate and optimize the conformer ensembles and **AIMNet-NSE** [29] is used to calculate three important quantum chemical descriptors: ionization potential, electron affinity, and electronegativity [30]. Note that **Auto3D** and **AIMNet-NSE** achieve DFT-level accuracy but are much more efficient [21, 31, 32].

- Ionization Potential (IP) is the minimum energy required to remove an electron from a neutral atom or molecule to form a positively charged ion (cation):  $IP = E_{\text{cation}} - E_{\text{neutral}}$ .
- Electron Affinity (EA) denotes the energy change associated with the addition of an electron to a neutral atom or molecule to form a negatively charged ion (anion):  $EA = E_{\text{neutral}} - E_{\text{anion}}$ .
- Electronegativity ( $\chi$ ) measures the tendency of an atom to attract a bonding pair of electrons:

$$\chi = -\left(\frac{\partial E}{\partial N}\right).$$

$E_{\text{cation}}$ ,  $E_{\text{neutral}}$ , and  $E_{\text{anion}}$  are the electronic energy of the positively charged, neutral, and negatively charged molecules, respectively.  $E$  and  $N$  are the energy and the number of electrons, respectively.

The tasks are to predict the Boltzmann-averaged value of each property across the conformer ensemble  $\langle y \rangle_{k_B} = \sum_{C_i \in \mathcal{C}} p_i y_i$ , where  $y_i$  is a conformer-specific property. We are given each  $C_i$ , and the goal is to predict  $\langle y \rangle_{k_B}$  from the molecular graph  $G$ , a single conformer  $C_i \in \mathcal{C}$ , or the set  $\mathcal{C}$ .

**Kraken** [33] is a dataset of 1,552 monodentate organophosphorus (III) ligands along with their DFT-computed conformer ensembles. In this study, we consider four 3D ligand descriptors exhibiting significant variance among conformers: Sterimol  $B_5$ , Sterimol L, buried Sterimol  $B_5$ , and buried Sterimol L. These descriptors quantify the steric features of each ligand in units of Å and are often employed for Quantitative Structure-Activity Relationship (QSAR) modeling in catalyst design.

As in the **Drugs-75K** tasks, the goal is to predict the Boltzmann-averaged value of each property across the conformer ensemble from the molecular graph  $G$ , a single conformer  $C_i \in \mathcal{C}$ , or the set  $\mathcal{C}$ .

Table 1: Statistics of the four datasets. The numbers of heavy atoms and rotatable bonds (“rot. bonds”) are averaged per conformer.

| Dataset   | # Molecules | # Conformers                      | # Heavy atoms  | # Rot. bonds | # Targets | Atomic species                                    |
|-----------|-------------|-----------------------------------|----------------|--------------|-----------|---------------------------------------------------|
| Drugs-75K | 75,099      | 558,002                           | 30.56          | 7.53         | 3         | H, C, N, O, F, Si, P, S, Cl                       |
| Kraken    | 1,552       | 21,287                            | 23.70          | 9.05         | 4         | H, B, C, N, O, F, Si, P, S, Cl, Fe, Se, Br, Sn, I |
| Dataset   | # Reactions | # Conformers                      | # Heavy atoms  | # Rot. bonds | # Targets | Atomic species                                    |
| EE        | 872         | Pro-R: 14,807<br>Pro-S: 13,999    | 59.32          | 18.57        | 1         | H, C, N, O, F, P, Cl, Br, Rh                      |
| BDE       | 5,915       | Ligand: 73,834<br>Complex: 40,264 | 29.62<br>32.38 | 6.99<br>6.99 | 1         | H, C, N, O, F, P, Cl, Ni, Cu, Br, Pd, Ag, Pt, Au  |

**EE** [34] is a dataset of 872 catalyst-substrate pairs involving 253 Rhodium (Rh)-bound atropisomeric catalysts derived from chiral bisphosphine, with 10 enamides as substrates. The dataset includes conformations of catalyst-substrate transition state complexes in two separate pro-S and pro-R configurations. The task is to predict the Enantiomeric Excess (EE) of the chemical reaction involving the substrate, defined as the absolute ratio between the concentration of each enantiomer in the product distribution. This dataset is generated with Q2MM, which automatically generates Transition State Force Fields (TSFFs) in order to simulate the conformer ensembles of each prochiral transition state complex. EE can then be computed from the conformer ensembles by Boltzmann-averaging the activation energies for the competing transition states [34, 35].

Unlike properties in Drugs-75K and Kraken, EE depends on the conformer ensembles of *each* pro-R and pro-S complex. The goal is to predict EE from the graphs of the catalyst and substrate ( $G_{\text{cat}}, G_{\text{sub}}$ ), a conformer  $C_i^{(R)} \in \mathcal{C}^{(R)}$  and  $C_i^{(S)} \in \mathcal{C}^{(S)}$  for each complex, or the ensembles  $\mathcal{C}^{(R)}$  and  $\mathcal{C}^{(S)}$ .

**BDE** [36] is a dataset containing 5,915 organometallic catalysts  $\text{ML}_1\text{L}_2$  consisting of a metal center ( $M = \text{Pd}, \text{Pt}, \text{Au}, \text{Ag}, \text{Cu}, \text{Ni}$ ) coordinated to two flexible organic ligands ( $\text{L}_1$  and  $\text{L}_2$ ), each selected from a 91-membered ligand library. The data includes conformations of each unbound catalyst, as well as conformations of the catalyst when bound to ethylene and bromide after oxidative addition with vinyl bromide. Each catalyst has an electronic binding energy, computed as the difference in the minimum energies of the bound-catalyst complex and unbound catalyst, following the DFT-optimization of their respective conformer ensembles. Although the binding energies are computed via DFT, the conformers provided for modeling are initially generated with Open Babel [37] followed by further geometry optimization, which ensures that the 3D structures are likely to be the global minimum energy conformers at the force field level [36]. Note that obtaining DFT-optimized conformers for BDE is not feasible given the significant computational cost. Therefore, this realistically represents the setting in which precise conformer ensembles are unknown at inference.

The task is to predict the binding energy from the graphs of the unbound and bound catalyst, sampled conformers  $C_i^{(\text{unbound})} \in \mathcal{C}^{(\text{unbound})}$  and  $C_i^{(\text{bound})} \in \mathcal{C}^{(\text{bound})}$ , or the ensembles  $\mathcal{C}^{(\text{unbound})}$  and  $\mathcal{C}^{(\text{bound})}$ .

**Dataset Preparation.** We implement several preprocessing steps to ensure the quality and validity of our datasets and facilitate their integration into machine learning models.

- **Conformer deduplication.** To eliminate redundant conformers in each ensemble  $\mathcal{C}$ , we first align every pair of conformers using RDKit [38], accounting for symmetric atom permutations. Subsequently, we employ Butina clustering [39] based on the Root Mean Square Deviation (RMSD) values derived from conformer alignment. Within each cluster, we select the conformer with the lowest energy. Note that Boltzmann-averaged regression labels are computed *before* deduplication.
- **Selection of molecules.** We focus on modeling flexible molecules, for which conformer ensemble learning may be relevant to capture their properties. Hence, we only retain molecules with more than 5 rotatable bonds. We also remove molecules with missing 3D geometries or 2D graphs.

## 4 BENCHMARKING MOLECULAR REPRESENTATION LEARNING MODELS

The representation of molecular data is crucial for applying machine learning models to problems in chemistry and biology. These representations typically include 1D strings, 2D topological graphs, and 3D geometric graphs. For a comprehensive benchmark for MRL models, our MARCEL includes a diverse representative selection of models for each of the aforementioned molecular representations.

In this section, we provide an overview of these models and describe how they are tailored to our tasks. We also introduce two strategies of explicitly encoding conformer ensembles using 3D models.

#### 4.1 1D MODELS

Our 1D baselines include Random Forest [40] models operating on molecular fingerprints [38, 41, 42]. Fingerprints convert a molecular graph into a bit array indicating the presence of chemical substructures and are widely used for cheminformatics and QSAR modeling in the low-data regime. Additionally, we include Long Short-Term Memory (LSTM) [43] and Transformer [44] models, popular sequence-based neural network architectures, operating on SMILES strings. For the BDE and EE datasets, we concatenate the SMILES of each molecule or complex with a “.” symbol and use a single sequence encoder. Further details on model implementations can be found in Appendix B.1.

#### 4.2 2D GRAPH NEURAL NETWORKS

We employ four widely-used GNN models as 2D baseline methods, including Graph Isomorphism Network (GIN) [45], GIN with Virtual Node (GIN-VN) [46], ChemProp [47], and GraphGPS [48]. GIN is a commonly-used model with strong representation ability. GIN-VN augments the vanilla GIN by incorporating a virtual node to aggregate the features of all nodes in the graph, thereby capturing global information more effectively. ChemProp is a directed message passing GNN designed specifically for molecular property prediction. GraphGPS is a Transformer-like [44] model specifically tailored for graph-structured data, which is able to capture long-range relationships.

Following OGB protocols [46], we employ a diverse set of atomic features such as aromaticity and hybridization for nodes, as well as bond features like ring information for edges (Appendix B.2). For the EE and BDE datasets, we employ a two-tower architecture with two separate 2D GNN models: for EE, since both pro-S and pro-R complexes share the same 2D graph, we leverage two separate GNNs to encode the catalyst and substrate; for BDE, we also encode the unbound and bound catalysts separately. We then concatenate these together to obtain the system-level embeddings.

#### 4.3 3D GRAPH NEURAL NETWORKS

We include six representative 3D GNNs that encompass diverse modeling perspectives. SchNet [12], an E(3)-invariant network, models spatial interactions by encoding pairwise interatomic distance. DimeNet++ [13], another E(3)-invariant model, uses directional message passing that embeds angles between triplets of atoms in order to enhance geometric expressivity. GemNet [14], an SE(3)-invariant model, utilizes a unique attention mechanism and dihedral angles between four atoms to model atomic interactions. PaiNN [15], initially developed to predict tensorial properties and molecular spectra, incorporates rotational equivariance into its message passing framework. ClofNet [16], an SE(3)-equivariant model that improves the popular EGNN [49], uses complete local frames for each atom, effectively capturing 3D atomistic structures while preserving invariance and equivariance. LEFTNet [17], based on ClofNet, introduces Local Substructure Encoding (LSE) and Frame Transition Encoding (FTE) to enhance the model expressivity via scalarization and tensorization.

We use atom types as the sole atom features for the 3D models. For both training and inference on Drug-75K, Kraken, and EE datasets, all the single-conformer 3D models encode the lowest-energy conformer of each conformer ensemble, which has the largest Boltzmann weight and hence provides the strongest model. Since imprecise conformers are encoded for the BDE task, we use a fixed, randomly sampled conformer for each unbound- and bound-catalyst during training and inference.

The 3D models also employ a two-tower architecture for the EE and BDE datasets. Two separate 3D GNNs are used to encode representations for each pro-S and pro-R complex in EE, or for each catalyst and bound complex in BDE, which are then concatenated to form the final representations.

We note that although using the lowest-energy conformer will yield the strongest performance, this setting can be unrealistic: it is often not possible to identify the lowest energy conformer without searching the entire conformer space. The lowest energy conformer can also depend on the force field used for geometry optimization, which may neglect experimental conditions such as solvents.

#### 4.4 INCORPORATING CONFORMER ENSEMBLES INTO MOLECULAR REPRESENTATIONS

3D geometric models primarily focus on learning representations from individual 3D structures. Although some models preserve global symmetries such as SE(3)-equivariance, these models do not learn representations that capture conformational flexibility which is caused by internal degrees of freedom such as bond rotations. Here, we describe two straightforward strategies that model conformational flexibility by explicitly leveraging conformer ensembles.

##### 4.4.1 STRATEGY 1: TRAINING-TIME DATA AUGMENTATION VIA CONFORMER SAMPLING

A direct approach to modeling conformer flexibility is to simply enrich the training data by randomly sampling a conformer from the ensemble during each training epoch. Formally, for a given molecule  $G$  and its conformer ensemble  $\mathcal{C}$ , we randomly select a conformer with uniform probability  $p = 1/|\mathcal{C}|$  while using the same training label for each conformer. Note that during inference, the lowest-energy conformer is used to evaluate the model performance. This strategy aligns with learning representations invariant to conformational changes, thus implicitly encoding the flexibility of molecular structures, and has been shown to be useful for learning chirality-sensitive 3D representations [19]. When conformer ensembles are available, the strategy is computationally efficient as it maintains the same complexity as the base 3D model. Unlike the other ensemble methods, this strategy can be used if conformer ensembles are only available at training time. In Appendix C, we evaluate two alternative scenarios where conformer ensembles are also available during evaluation.

##### 4.4.2 STRATEGY 2: ENSEMBLE LEARNING WITH EXPLICIT SET ENCODERS

The second strategy utilizes a set encoder to simultaneously encode the entire conformer ensemble  $\mathcal{C}$  at both training and inference time. Inspired by the multi-instance learning framework [50–52], this strategy first employs 3D GNNs to generate individual conformer embeddings and then aggregates these embeddings using a set encoder, as illustrated in Figure 2.

Formally, for each conformer  $C_i \in \mathcal{C}$ , we obtain its corresponding embedding  $z_i = f(G, C_i) \in \mathbb{R}^d$ , where  $f$  is a single-conformer 3D model and  $d$  is the embedding dimension. Note that the embedding  $z$  is a (3D) graph-level representation resulting from a pooling function over the node-level embeddings after message passing. To further aggregate these embeddings  $\mathcal{Z} = \{z_i\}_{i=1}^{|\mathcal{C}|}$  into a single molecular representation, we consider the following three set encoders:

- **Mean pooling** simply computes the mean of all the conformer embeddings:

$$s^{\text{MEAN}} = \frac{1}{|\mathcal{C}|} \sum_{i=1}^{|\mathcal{C}|} z_i. \quad (1)$$

- **DeepSets** [53] utilizes a permutation-invariant function to process a set of inputs. It first applies a MultiLayer Perceptron (MLP)  $h$  to each conformer embedding and then aggregates the transformed embeddings using sum pooling followed by another MLP  $g$ :

$$s^{\text{DS}} = g \left( \sum_{i=1}^{|\mathcal{C}|} h(z_i) \right). \quad (2)$$

This method retains more discernible information from individual embeddings compared to mean pooling at a cost of two non-linear functions.

- **Self-attention** [54] further computes a weighted sum of the embeddings, where the weights are obtained by applying a softmax function to the dot product of the embeddings:

$$s^{\text{ATT}} = \sum_{i=1}^{|\mathcal{C}|} c_i, \quad \text{where } c_i = g \left( \sum_{j=1}^{|\mathcal{C}|} \alpha_{ij} h(z_j) \right), \quad \alpha_{ij} = \frac{\exp((Wh(z_i))^T (Wh(z_j)))}{\sum_{k=1}^{|\mathcal{C}|} \exp((Wh(z_i))^T (Wh(z_k)))}. \quad (3)$$

Here,  $W \in \mathbb{R}^{d \times d}$  is a learnable weight matrix. This approach can capture conformer interactions.

By employing these set encoders, we can learn a model that is more sensitive to the full range of conformer variations present in the ensemble. After obtaining the ensemble embeddings, we further apply a linear projection head to generate the final prediction.

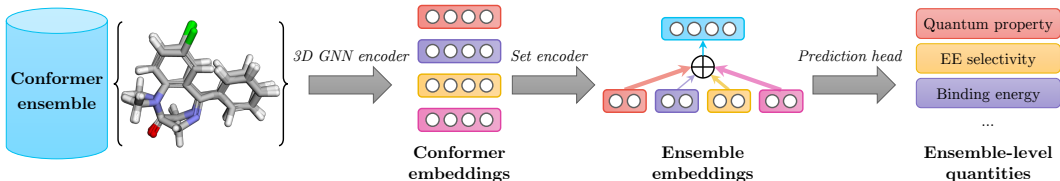


Figure 2: Conformer ensemble learning with explicit set encoders (Strategy 2). Individual conformer embeddings are first obtained via 3D GNN encoders. Then, a set encoder is employed to aggregate conformer embeddings. Finally, a linear projection head is used to generate the prediction.

## 5 EXPERIMENTS

### 5.1 EXPERIMENTAL CONFIGURATIONS

Each dataset is partitioned randomly into three subsets: 70% for training, 10% for validation, and 20% for test. Each model is trained over 2,000 epochs using the Adam optimizer [55] with early stopping triggered if there is no improvement on the training loss over 200 epochs. For all nine regression targets, experiments are repeated three times, and the results reported correspond to the model that performs best on the validation set in terms of Mean Absolute Error (MAE).

The Boltzmann-averaged targets are computed over all available conformers. For ensemble learning models, we cap the number of encoded conformers per molecule to a maximum of 20, which empirically improves training stability and leads to better performance. To ensure a fair comparison, the hidden dimension size is uniformly set to 128 for all models. Other settings mostly follow the original configurations as described in the respective papers. We specify all hyperparameters and describe experimental environments in Appendix B.3.

### 5.2 RESULTS AND ANALYSIS

We summarize the performance of the 1D, 2D, and 3D MRL models and the best results from ensemble learning strategies on 3D models in Table 2. Figure 3 reports the *performance changes* in Mean Absolute Error (MAE) for each 3D model when applying the ensemble learning strategies. The raw performance data with standard deviation and the parameter size of each model can be found in Appendix D. In summary, although performance varies across the datasets, tasks, and models, the ensemble learning strategies improve upon 3D models that only encode one conformer in 48 out of 54 experiments across 9 tasks and 6 base models, demonstrating the effectiveness of conformer ensemble learning. Our analysis leads to the following key observations.

**Observation 1. The conformer ensemble learning strategy with explicit set encoders frequently yields improved performance.**

Figure 3 indicates that encoding conformer ensembles can substantially reduce test error, achieving improvements in 108 experiments across all 9 tasks, 6 base models, and 3 set encoders, most notably on the tasks in the smaller-sized Kraken dataset. This, however, does not always extend to larger datasets like Drugs-75K. We conjecture that for Drugs-75K, the computational burden of encoding all conformers in each ensemble alters the learning dynamics of the underlying model, making training more challenging. A similar finding was reported by Axelrod and Gómez-Bombarelli [23].

Among the three set encoders, DeepSets consistently demonstrates significant improvements in 42 out of 54 experiments across 9 tasks and 6 base 3D models. We conjecture that this superior performance is due to its ability of effectively modeling set objects at a relatively minor computational overhead of two non-linear transformations. On the other hand, the simple mean pooling approach loses discriminative power across the conformers in the ensemble, resulting in inferior performance. It is also noteworthy that the attention models exhibit mixed results compared to the vanilla 3D models, despite theoretically being the most powerful set encoders. This inconsistency might be attributable to the computational intricacy of the self-attention layer, which models the pairwise relationship among conformers in each ensemble and hence could require more sophisticated training strategies. Future research should consider developing better neural architectures that are specifically designed to more efficiently encode structural information from conformer ensembles.

Table 2: Performance of 1D, 2D, and 3D baseline MRL models and the best results from ensemble learning strategies on 3D GNNs. The metric used is the Mean Absolute Error (MAE,  $\downarrow$ ). The **bold** indicates the best-performing model, while underlined denotes the second-best.

| Category               | Model         | Drugs-75K     |               |               | Kraken         |               |                   |               | EE             | BDE           |
|------------------------|---------------|---------------|---------------|---------------|----------------|---------------|-------------------|---------------|----------------|---------------|
|                        |               | IP            | EA            | $\chi$        | B <sub>5</sub> | L             | BurB <sub>5</sub> | BurL          |                |               |
| 1D                     | Random forest | 0.4987        | 0.4747        | 0.2732        | 0.4760         | 0.4303        | 0.2758            | 0.1521        | 61.2963        | 3.0335        |
|                        | LSTM          | 0.4788        | 0.4648        | 0.2505        | 0.4879         | 0.5142        | 0.2813            | 0.1924        | 64.0088        | 2.8279        |
|                        | Transformer   | 0.6617        | 0.5850        | 0.4073        | 0.9611         | 0.8389        | 0.4929            | 0.2781        | 62.0816        | 10.0771       |
| 2D                     | GIN           | 0.4354        | 0.4169        | 0.2260        | 0.3128         | 0.4003        | 0.1719            | 0.1200        | 62.3065        | 2.6368        |
|                        | GIN+VN        | 0.4361        | 0.4169        | 0.2267        | 0.3567         | 0.4344        | 0.2422            | 0.1741        | 62.3815        | 2.7417        |
|                        | ChemProp      | 0.4595        | 0.4417        | 0.2441        | 0.4850         | 0.5452        | 0.3002            | 0.1948        | 61.0336        | 2.6616        |
|                        | GraphGPS      | 0.4351        | 0.4085        | 0.2212        | 0.3450         | 0.4363        | 0.2066            | 0.1500        | 61.6251        | 2.4827        |
| 3D                     | SchNet        | 0.4394        | 0.4207        | 0.2243        | 0.3293         | 0.5458        | 0.2295            | 0.1861        | 17.7421        | 2.5488        |
|                        | DimeNet++     | 0.4441        | 0.4233        | 0.2436        | 0.3510         | 0.4174        | 0.2097            | 0.1526        | 14.6414        | <b>1.4503</b> |
|                        | GemNet        | <u>0.4069</u> | <u>0.3922</u> | <b>0.1970</b> | 0.2789         | 0.3754        | 0.1782            | 0.1635        | 18.0338        | 1.6530        |
|                        | PaiNN         | 0.4505        | 0.4495        | 0.2324        | 0.3443         | 0.4471        | 0.2395            | 0.1673        | 20.2359        | 2.1261        |
|                        | ClofNet       | 0.4393        | 0.4251        | 0.2378        | 0.4873         | 0.6417        | 0.2884            | 0.2529        | 33.9473        | 2.6057        |
|                        | LEFTNet       | 0.4174        | 0.3964        | 0.2083        | 0.3072         | 0.4493        | 0.2176            | 0.1486        | 19.7974        | 1.5328        |
| Best Ensemble Strategy | SchNet        | 0.4452        | 0.4232        | 0.2243        | 0.2704         | 0.4322        | 0.2024            | 0.1443        | 14.2238        | 1.9737        |
|                        | DimeNet++     | 0.4126        | 0.3944        | 0.2267        | 0.2630         | 0.3468        | 0.1783            | 0.1185        | 12.0259        | 1.4741        |
|                        | GemNet        | <b>0.4066</b> | <b>0.3910</b> | 0.2027        | 0.2313         | <b>0.3386</b> | <b>0.1589</b>     | <b>0.0947</b> | <b>11.6142</b> | 1.6059        |
|                        | PaiNN         | 0.4466        | 0.4269        | 0.2294        | <b>0.2225</b>  | 0.3619        | 0.1693            | 0.1324        | 13.5570        | 1.8744        |
|                        | ClofNet       | 0.4280        | 0.4033        | 0.2199        | 0.3228         | 0.4485        | 0.2178            | 0.1548        | 13.9647        | 2.0106        |
|                        | LEFTNet       | 0.4149        | 0.3953        | 0.2069        | 0.2644         | 0.3643        | 0.2017            | 0.1386        | 18.4189        | 1.5276        |

**Observation 2. Sampling conformers at training time can improve performance, especially on imprecise conformer structures.**

We observe that data augmentation improves performance on 34 experiments, especially on the challenging BDE dataset, where the other ensemble learning strategies often do not help. Note that the conformers in the BDE dataset originate from Open Babel, as opposed to the golden-standard DFT-level conformers present in other datasets. This suggests that training with randomly sampled conformers might offer robustness to noise in the imprecise structures. On other tasks, randomly sampling the conformers at each epoch may help the model learn an invariance to conformational changes, but does not always increase performance for all 3D models. This might be because the sampling probability is uniform across the entire conformer set, which does not respect the underlying Boltzmann weight of each conformer. In future work, it may be worthwhile to investigate whether more physics-informed sampling strategies could lead to more consistent performance gains.

**Observation 3. No model consistently outperforms the rest, with substantial task dependencies.**

The results in Table 2 suggest that no single model outperforms the others across all tasks. Of the 1D models, LSTM outperforms Random Forest and Transformer models on Drugs-75K and BDE, demonstrating the effectiveness of SMILES-based representations of molecules on large-scale datasets. For small datasets such as Kraken and EE, Random Forests outperform sequence models at a lower computational cost, indicating that traditional models are superior in the low-data regime.

Amongst 2D models, GIN delivers the best performance on four tasks compared to all other models; GraphGPS also demonstrates strong performance on several tasks (B<sub>5</sub>, L, and BurL). Surprisingly, the 2D models are also competitive with some 3D models on the large-scale Drugs-75K tasks. This is possibly due to the fact that the electronic properties in Drugs-75K are not as sensitive to conformational changes, thus explicitly modeling the structures in 3D may not be necessary. However, all 2D models perform worse as compared to the 3D models in the reaction datasets EE and BDE, indicating the important role of spatial interactions in determining reaction-related properties.

For 3D models, GemNet and LEFTNet excel in IP, EA, and  $\chi$ . The complexity of these two equivariant models may especially benefit from the large dataset size of Drugs-75K. For Kraken and the two reaction datasets, DimeNet++ — an invariant model — achieves promising performance, suggesting that highly-complex 3D models may be less useful for chemical applications with small-to-medium sized datasets. On EE, we observe that 3D models remarkably outperform 1D and 2D models, likely because enantioselectivity depends on subtle spatial interactions. When predicting binding energies, using 3D models also leads to modest improvements.

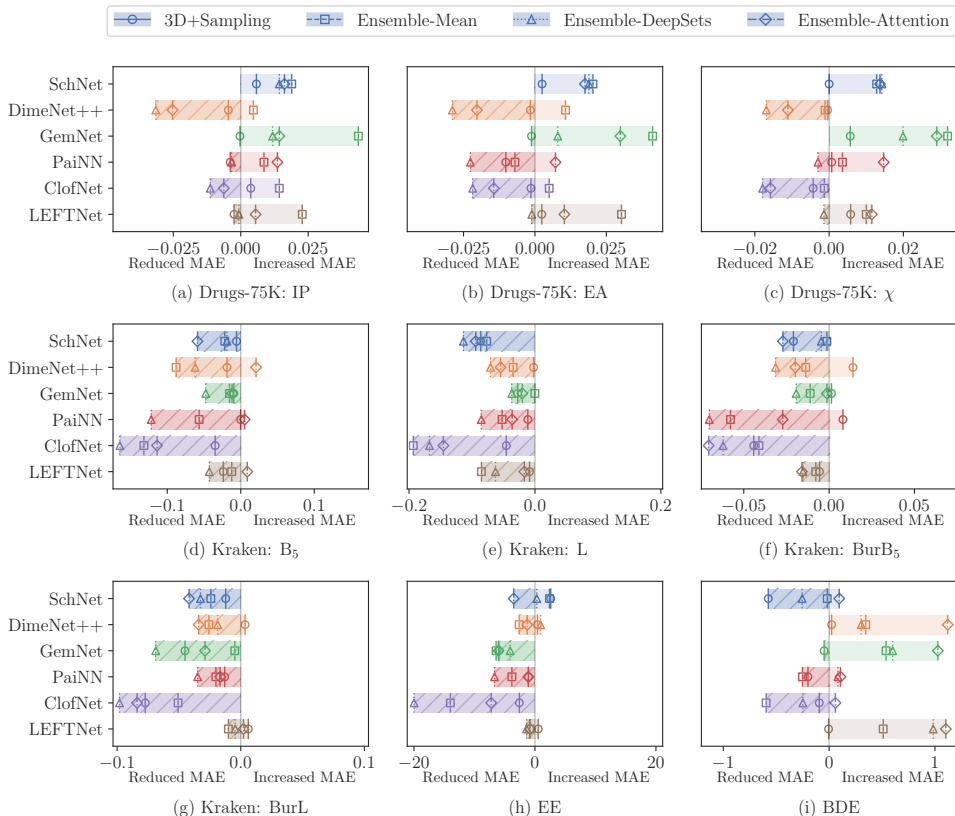


Figure 3: *Performance changes* of four conformer ensemble learning strategies on the basis of six 3D graph models. Here, negative values (marked in hatch patterns) denote *reduced* Mean Absolute Error (MAE), signifying a performance improvement due to the incorporation of conformer ensembles.

Overall, model performance varies substantially across tasks, even within the same dataset, emphasizing the diversity of the tasks in MARCEL. Generally, 1D and 2D models perform well on small-scale molecular datasets, while 3D models excel on large datasets and reaction-centric tasks. MARCEL also highlights the benefits of explicitly encoding multiple conformers to improve MRL.

## 6 DISCUSSIONS AND CONCLUSIONS

In this work, we present the first Molecular Conformer Ensemble Learning benchmark (MARCEL) to evaluate the potential of learning from a set of conformer structures. Through two conformer ensemble learning strategies, we discover performance improvements across various tasks. However, there are some limitations that require further consideration. First, our studied ensemble learning strategies do not universally improve performance across all tasks and datasets. This highlights the need for more tailored approaches that integrate with domain expertise to better model specific tasks and datasets of practical interest. Second, the computational cost of encoding all conformers within the ensembles, especially for larger datasets, suggests the need to further study the trade-offs between model complexity and efficiency. Finally, our datasets only contain regression tasks and do not cover all of the relevant chemical space, which might limit the generalization of our experimental findings.

Despite these challenges, we envision that our work will prompt further research in the geometric deep learning community on how to make use of conformer ensembles for molecular property prediction. For instance, future research could explore new model architectures that can efficiently encode ensemble-level information or more sophisticated conformer sampling strategies. We also hope that our work will stimulate collaborative research across the machine learning and chemistry fields, with the ultimate goal of pushing the boundaries of predictive molecular modeling and aligning algorithmic advancements with the practical needs of the chemistry community.

## ACKNOWLEDGEMENTS

This work is supported by NSF Center for Computer Assisted Synthesis (2202693).

## REFERENCES

- [1] Jun Xia, Yanqiao Zhu, Yuanqi Du, and Stan Z. Li. A Systematic Survey of Chemical Pre-trained Models. 2023. [1](#)
- [2] Oliver Wieder, Stefan Kohlbacher, Méline Kuenemann, Arthur Garon, Pierre Ducrot, Thomas Seidel, and Thierry Langer. A Compact Review of Molecular Property Prediction with Graph Neural Networks. *Drug Discov. Today Technol.*, 37:1–12, 2020. [1](#)
- [3] W. Patrick Walters and Regina Barzilay. Applications of Deep Learning in Molecule Generation and Molecular Property Prediction. *Acc. Chem. Res.*, 54(2):263–270, 2021. [1](#)
- [4] H. L. Morgan. The Generation of a Unique Machine Description for Chemical Structures — A Technique Developed at Chemical Abstracts Service. *J. Chem. Doc.*, 5(2):107–113, 1965. [1](#)
- [5] Robert C. Glem, Andreas Bender, Catrin H. Arnby, Lars Carlsson, Scott Boyer, and James Smith. Circular Fingerprints: Flexible Molecular Descriptors with Applications from Physical Chemistry to ADME. *IDrugs*, 9(3):199–204, 2006. [1](#)
- [6] G. Skoraczynski, P. Dittwald, B. Miasojedow, S. Szymkuć, E. P. Gajewska, B. A. Grzybowski, and A. Gambin. Predicting the Outcomes of Organic Reactions via Machine Learning: Are Current Descriptors Sufficient? *Sci. Rep.*, 7(1):1–9, 2017. [1](#)
- [7] Zhen Liu, Yurii S. Moroz, and Olexandr Isayev. The Challenge of Balancing Model Sensitivity and Robustness in Predicting Yields: A Benchmarking Study of Amide Coupling Reactions. *chemrxiv.org*, 2023. [1](#)
- [8] Thomas N. Kipf and Max Welling. Semi-Supervised Classification with Graph Convolutional Networks. In *ICLR*, 2017. [1](#)
- [9] Justin Gilmer, Samuel S. Schoenholz, Patrick F. Riley, Oriol Vinyals, and George E. Dahl. Neural Message Passing for Quantum Chemistry. In *ICML*, pages 1263–1272, 2017. [1](#)
- [10] Keyulu Xu, Chengtao Li, Yonglong Tian, Tomohiro Sonobe, Ken-ichi Kawarabayashi, and Stefanie Jegelka. Representation Learning on Graphs with Jumping Knowledge Networks. In *ICML*, pages 5453–5462, 2018. [1](#)
- [11] Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Liò, and Yoshua Bengio. Graph Attention Networks. In *ICLR*, 2018. [1](#)
- [12] Kristof Schütt, Pieter-Jan Kindermans, Huziel Enoc Sauceda Felix, Stefan Chmiela, Alexandre Tkatchenko, and Klaus-Robert Müller. SchNet: A Continuous-Filter Convolutional Neural Network for Modeling Quantum Interactions. In *NIPS*, pages 991–1001, 2017. [1](#), [5](#)
- [13] Johannes Klicpera, Janek Groß, and Stephan Günnemann. Directional Message Passing for Molecular Graphs. In *ICLR*, 2020. [1](#), [5](#)
- [14] Johannes Gasteiger, Florian Becker, and Stephan Günnemann. GemNet: Universal Directional Graph Neural Networks for Molecules. In *NeurIPS*, pages 6790–6802, 2021. [1](#), [5](#)
- [15] Kristof Schütt, Oliver T. Unke, and Michael Gastegger. Equivariant Message Passing for the Prediction of Tensorial Properties and Molecular Spectra. In *ICML*, pages 9377–9388, 2021. [1](#), [5](#)
- [16] Weitao Du, He Zhang, Yuanqi Du, Qi Meng, Wei Chen, Nanning Zheng, Bin Shao, and Tie-Yan Liu. SE(3) Equivariant Graph Neural Networks with Complete Local Frames. In *ICML*, pages 5583–5608, 2022. [1](#), [5](#)
- [17] Weitao Du, Yuanqi Du, Limei Wang, Dieqiao Feng, Guifeng Wang, Shuiwang Ji, Carla Gomes, and Zhi-Ming Ma. A New Perspective on Building Efficient and Expressive 3D Equivariant Graph Neural Networks. *arXiv.org*, 2023. [1](#), [5](#)
- [18] Christopher Morris, Martin Ritzert, Matthias Fey, William L. Hamilton, Jan Eric Lenssen, Gaurav Rattan, and Martin Grohe. Weisfeiler and Leman Go Neural: Higher-Order Graph Neural Networks. In *AAAI*, pages 4602–4609, 2019. [1](#)

- [19] Keir Adams, Lagnajit Pattanaik, and Connor W. Coley. Learning 3D Representations of Molecular Chirality with Invariance to Bond Rotations. In *ICLR*, 2022. 1, 6
- [20] Bharath Ramsundar, Peter Eastman, Patrick Walters, and Vijay Pande. *Deep Learning for the Life Sciences: Applying Deep Learning to Genomics, Microscopy, Drug Discovery, and More*. O'Reilly Media, 2019. 2
- [21] Emanuele Perola and Paul S. Charifson. Conformational Analysis of Drug-Like Molecules Bound to Proteins: An Extensive Study of Ligand Reorganization upon Binding. *J. Med. Chem.*, 47(10):2499–2510, 2004. 2, 3
- [22] Andrew F. Zahrt, Jeremy J. Henle, Brennan T. Rose, Yang Wang, William T. Darrow, and Scott E. Denmark. Prediction of Higher-Selectivity Catalysts by Computer-Driven Workflow and Machine Learning. *Science*, 363(6424):eaau5631, 2019. 2
- [23] Simon Axelrod and Rafael Gómez-Bombarelli. Molecular Machine Learning with Conformer Ensembles. *arXiv.org*, 2020. 2, 7
- [24] Jan Weinreich, Nicholas J. Browning, and O. Anatole von Lilienfeld. Machine Learning of Free Energies in Chemical Compound Space Using Ensemble Representations: Reaching Experimental Uncertainty for Solvation. *J. Chem. Phys.*, 154(13):134113, 2021. 2
- [25] Kangway V. Chuang and Michael J. Keiser. Attention-Based Learning on Molecular Ensembles. *arXiv.org*, 2020. 2
- [26] Chaitanya K. Joshi, Arian R. Jamasb, Ramon Viñas, Charles Harris, Simon Mathis, and Pietro Liò. Multi-State RNA Design with Geometric Multi-Graph Neural Networks. *arXiv.org*, 2023. 2
- [27] Simon Axelrod and Rafael Gómez-Bombarelli. GEOM, Energy-Annotated Molecular Conformations for Property Prediction and Molecular Generation. *Sci. Data*, 9(1):185, 2022. 3, 14
- [28] Zhen Liu, Tetiana Zubatiuk, Adrian Roitberg, and Olexandr Isayev. Auto3D: Automatic Generation of the Low-Energy 3D Structures with ANI Neural Network Potentials. *J. Chem. Inf. Model.*, 62(22):5373–5382, 2022. 3, 14
- [29] Roman Zubatyuk, Justin S. Smith, Benjamin T. Nebgen, Sergei Tretiak, and Olexandr Isayev. Teaching a Neural Network to Attach and Detach Electrons From Molecules. *Nat. Commun.*, 12(1):1–11, 2021. 3, 14
- [30] Andrzej M. Żurański, Jason Y. Wang, Benjamin J. Shields, and Abigail G. Doyle. Auto-QChem: An Automated Workflow for the Generation and Storage of DFT Calculations for Organic Molecules. *React. Chem. Eng.*, 7(6):1276–1284, 2022. 3, 14
- [31] Qiyuan Zhao, Sai Mahit Vaddadi, Michael Woulfe, Lawal A. Ogunfowora, Sanjay S. Garimella, Olexandr Isayev, and Brett M. Savoie. Comprehensive Exploration of Graphically Defined Reaction Spaces. *Sci. Data*, 10(1):1–10, 2023. 3, 14
- [32] Peikun Zheng, Roman Zubatyuk, Wei Wu, Olexandr Isayev, and Pavlo O. Dral. Artificial Intelligence-Enhanced Quantum Chemical Method with Broad Applicability. *Nat. Commun.*, 12(1):7022, 2021. 3
- [33] Tobias Gensch, Gabriel dos Passos Gomes, Pascal Friederich, Ellyn Peters, Théophile Gaudin, Robert Pollice, Kjell Jorner, AkshatKumar Nigam, Michael Lindner-D’Addario, Matthew S. Sigman, and Alán Aspuru-Guzik. A Comprehensive Discovery Platform for Organophosphorus Ligands for Catalysis. *J. Am. Chem. Soc.*, 144:1205–1217, 2022. 3, 15
- [34] Anthony R. Rosales, Jessica Wahlers, Elaine Limé, Rebecca E. Meadows, Kevin W. Leslie, Rhona Savin, Fiona Bell, Eric Hansen, Paul Helquist, Rachel H. Munday, Olaf Wiest, and Per-Ola Norrby. Rapid Virtual Screening of Enantioselective Catalysts Using CatVS. *Nat. Catal.*, 2(1):41–45, 2019. 4, 15, 16
- [35] G. P. Moss. Basic Terminology of Stereochemistry (IUPAC Recommendations 1996). *Pure Appl. Chem.*, 68(12):2193–2222, 1996. 4, 16
- [36] Benjamin Meyer, Boodsarin Sawatlon, Stefan Heinen, O. Anatole von Lilienfeld, and Clémence Corminboeuf. Machine Learning Meets Volcano Plots: Computational Discovery of Cross-Coupling Catalysts. *Chem. Sci.*, 9:7069–7077, 2018. 4, 16
- [37] Noel M. O’Boyle, Michael Banck, Craig A. James, Chris Morley, Tim Vandermeersch, and Geoffrey R. Hutchison. Open Babel: An Open Chemical Toolbox. *J. Cheminformatics*, 3(1):1–14, 2011. 4, 16

- [38] Greg Landrum, Paolo Tosco, Brian Kelley, Ric, sriniker, gedec, Riccardo Vianello, NadineSchneider, Eisuke Kawashima, Andrew Dalke, Dan N, David Cosgrove, Brian Cole, Matt Swain, Samo Turk, AlexanderSavelyev, Gareth Jones, Alain Vaucher, Maciej Wójcikowski, Ichiru Take, Daniel Probst, Kazuya Ujihara, Vincent F. Scalfani, guillaume godin, Axel Pahl, Francois Berenger, JLVarjo, strets123, JP, and DoliathGavid. rdkit/rdkit: 2022\_03\_2 (q1 2022) release, 2022. 4, 5, 16
- [39] Darko Butina. Unsupervised Data Base Clustering Based on Daylight’s Fingerprint and Tanimoto Similarity: A Fast and Automated Way To Cluster Small and Large Data Sets. *J. Chem. Inf. Comput. Sci.*, 39(4): 747–750, 1999. 4
- [40] Leo Breiman. Random Forests. *Mach. Learn.*, 45(1):5–32, 2001. 5
- [41] David Rogers and Mathew Hahn. Extended-Connectivity Fingerprints. *J. Chem. Inf. Model.*, 50:742–754, 2010. 5, 16
- [42] Joseph L. Durant, Burton A. Leland, Douglas R. Henry, and James G. Nourse. Reoptimization of MDL Keys for Use in Drug Discovery. *J. Chem. Inf. Comput. Sci.*, 42(6):1273–1280, 2002. 5, 16
- [43] Sepp Hochreiter and Jürgen Schmidhuber. Long Short-Term Memory. *Neural Comp.*, 9(8):1735–1780, 1997. 5
- [44] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Uszkoreit Kaiser, and Illia Polosukhin. Attention is All You Need. In *NIPS*, pages 5998–6008, 2017. 5, 16
- [45] Keyulu Xu, Weihua Hu, Jure Leskovec, and Stefanie Jegelka. How Powerful Are Graph Neural Networks? In *ICLR*, 2019. 5
- [46] Weihua Hu, Matthias Fey, Marinka Zitnik, Yuxiao Dong, Hongyu Ren, Bowen Liu, Michele Catasta, and Jure Leskovec. Open Graph Benchmark: Datasets for Machine Learning on Graphs. In *NeurIPS*, pages 22118–22133, 2020. 5, 17
- [47] Kevin Yang, Kyle Swanson, Wengong Jin, Connor W. Coley, Philipp Eiden, Hua Gao, Angel Guzman-Perez, Timothy Hopper, Brian Kelley, Miriam Mathea, Andrew Palmer, Volker Settels, Tommi S. Jaakkola, Klavs F. Jensen, and Regina Barzilay. Analyzing Learned Molecular Representations for Property Prediction. *J. Chem. Inf. Model.*, 59(8):3370–3388, 2019. 5
- [48] Ladislav Rampášek, Mikhail Galkin, Vijay Prakash Dwivedi, Anh Tuan Luu, Guy Wolf, and Dominique Beaini. Recipe for a General, Powerful, Scalable Graph Transformer. In *NeurIPS*, pages 14501–14515, 2022. 5
- [49] Victor Garcia Satorras, Emiel Hoogeboom, and Max Welling. E(n) Equivariant Graph Neural Networks. In *ICML*, pages 9323–9332, 2021. 5
- [50] Thomas G. Dietterich, Richard H. Lathrop, and Tomás Lozano-Pérez. Solving the Multiple Instance Problem with Axis-Parallel Rectangles. *Artif. Intell.*, 89(1-2):31–71, 1997. 6
- [51] Oded Maron and Tomás Lozano-Pérez. A Framework for Multiple-Instance Learning. In *NIPS*, pages 570–576, 1997. 6
- [52] Maximilian Ilse, Jakub M. Tomczak, and Max Welling. Attention-based Deep Multiple Instance Learning. In *ICML*, pages 2132–2141, 2018. 6
- [53] Manzil Zaheer, Satwik Kottur, Siamak Ravanbakhsh, Barnabás Póczos, Ruslan R. Salakhutdinov, and Alexander J. Smola. Deep Sets. In *NIPS*, pages 3391–3401, 2017. 6
- [54] Dzmitry Bahdanau, Kyunghyun Cho, and Yoshua Bengio. Neural Machine Translation by Jointly Learning to Align and Translate. In *ICLR*, 2015. 6
- [55] Diederik P. Kingma and Jimmy Ba. Adam: A Method for Stochastic Optimization. In *ICLR*, 2015. 7, 17
- [56] Carlo Adamo and Vincenzo Barone. Toward Reliable Density Functional Methods Without Adjustable Parameters: The PBE0 Model. *J. Chem. Phys.*, 110(13):6158–6170, 1999. 14
- [57] A. Verloop, W. Hoogenstraaten, and J. Tipker. Development and Application of New Steric Substituent Parameters in Drug Design. In E. J. Ariëns, editor, *Drug Design*, volume 11 of *Medicinal Chemistry: A Series of Monographs*, pages 165–207. Academic Press, Amsterdam, 1976. 15

- [58] Fabian Pedregosa, Gaël Varoquaux, Alexandre Gramfort, Vincent Michel, Bertrand Thirion, Olivier Grisel, Mathieu Blondel, Peter Prettenhofer, Ron Weiss, Vincent Dubourg, Jake VanderPlas, Alexandre Passos, David Cournapeau, Matthieu Brucher, Matthieu Perrot, and Edouard Duchesnay. Scikit-learn: Machine Learning in Python. *J. Mach. Learn. Res.*, 12:2825–2830, 2011. [16](#)
- [59] Philip Gage. A New Algorithm for Data Compression. *C Users J.*, 12(2):23–38, 1994. [16](#)
- [60] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, Alban Desmaison, Andreas Kopf, Edward Yang, Zachary DeVito, Martin Raison, Alykhan Tejani, Sasank Chilamkurthy, Benoit Steiner, Lu Fang, Junjie Bai, and Soumith Chintala. PyTorch: An Imperative Style, High-Performance Deep Learning Library. In *NeurIPS*, pages 8024–8035, 2019. [17](#)
- [61] Matthias Fey and Jan Eric Lenssen. Fast Graph Representation Learning with PyTorch Geometric. In *RLGM@ICLR*, 2019. [17](#)

## Supplementary Material for MARCEL

|                                                                                          |           |
|------------------------------------------------------------------------------------------|-----------|
| <b>A Dataset Description</b>                                                             | <b>14</b> |
| A.1 Drugs-75K . . . . .                                                                  | 14        |
| A.2 Kraken . . . . .                                                                     | 15        |
| A.3 EE . . . . .                                                                         | 15        |
| A.4 BDE . . . . .                                                                        | 16        |
| <b>B Implementation Details</b>                                                          | <b>16</b> |
| B.1 Implementation of 1D Models . . . . .                                                | 16        |
| B.2 Featurizations of Molecules for 2D Models . . . . .                                  | 17        |
| B.3 Hyperparameter Specifications and Experimental Environments . . . . .                | 17        |
| <b>C Additional Experiments on Evaluation Schemes of the Conformer Sampling Strategy</b> | <b>17</b> |
| <b>D Raw Data</b>                                                                        | <b>18</b> |

### A DATASET DESCRIPTION

MARCEL include four datasets that cover a diverse range of chemical space, which focuses on four chemically-relevant tasks for both molecules and reactions, with an emphasis on Boltzmann-averaged properties of conformer ensembles computed at the Density-Functional Theory (DFT) level. Detailed information regarding dataset access, data formatting, and loading procedures can be found at our GitHub repository <https://github.com/SXKDZ/MARCEL>. Any subsequent updates will also be posted on this repository.

#### A.1 DRUGS-75K

Drugs-75K is a subset of the GEOM-Drugs [27] dataset, which includes 75,099 drug-like molecules with at least 5 rotatable bonds. The original GEOM-Drugs dataset was constructed using semi-empirical DFT methods, which is less accurate than full DFT. To curate the Drugs-75K subset, Auto3D [28] is used to generate and optimize the conformer ensembles for each molecule and AIMNet-NSE [29] is used to calculate three important DFT-based reactivity descriptors: ionization potential, electron affinity, and electronegativity [30].

Auto3D [28] efficiently generates high-quality conformers, with a mean RMSD at around 0.2 Å when compared with DFT conformers. It has been used in other large conformer dataset generation [31]. Regarding the neural network surrogate AIMNET-NSE [29], it mimics the PBE0/ma-def2-SVP method of DFT, which is widely used in the chemistry community. Investigating their accuracy is out of the scope of this paper, but are readily accessible from multiple sources [29, 56].

**Objectives.** The tasks are to predict the Boltzmann-averaged value of each property across the conformer ensemble  $\langle y \rangle_{k_B} = \sum_{C_i \in \mathcal{C}} p_{C_i} y_{C_i}$ , where  $y_{C_i}$  is a conformer-specific property. We are given each  $C_i$ , and the goal is to predict  $\langle y \rangle_{k_B}$  from the molecular graph  $G$ , a single conformer  $C_i \in \mathcal{C}$ , or the set  $\mathcal{C}$ .

**Dataset preparation.** In preparing the 75K version of GEOM-Drugs, we begin with the original SMILES strings of the molecules. We first exclude molecules that have less than 5 rotatable bonds. To enable the utilization of AIMNet-NSE for descriptor computation, we retain only those molecules containing atoms of H, C, N, O, F, Si, P, S, and Cl. Further, we generate DFT-level conformers and compute their energies with Auto3D. Based on these conformers, we compute three chemical bond energy descriptors using AIMNet-NSE. We exclude conformers that Auto3D fails to converge and charged molecules that are unable to be processed by AIMNet-NSE, which results in 75,099

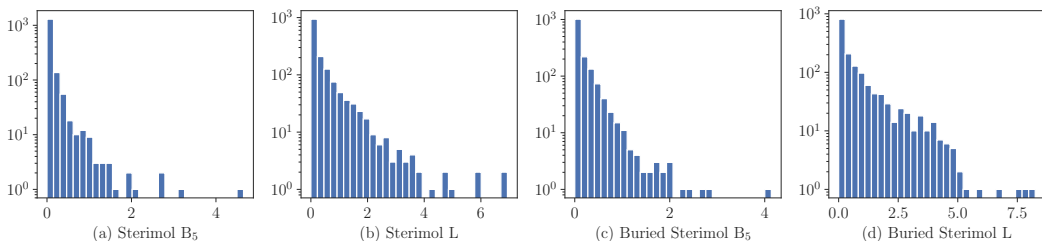


Figure S1: Histogram of the ratio of the variance of each conformer property to the variance of each Boltzmann-averaged property in the Kraken dataset.

molecules. Subsequently, we compute molecular-level Boltzmann-averaged descriptors based on conformer-level descriptors. Finally, we undertake a deduplication process as outlined in Section 3 with a RMSD threshold of 2.0, which yields a total of 558,002 distinct conformers.

**Data availability and license.** The original GEOM-Drugs dataset is publicly available at <https://github.com/learningmatter-mit/geom> but no license is specified. Our Drugs-75K can be accessed at <https://github.com/SXKDZ/MARCEL/tree/main/datasets/Drugs>. As for the conformer ensembles and descriptors that we generated, they are licensed under the Apache License.

## A.2 KRAKEN

Kraken [33] is a dataset of 1,552 monodentate organophosphorus (III) ligands along with their DFT-computed conformer ensembles. In this study, we consider four 3D catalytic ligand descriptors exhibiting significant variance among conformers: Sterimol B<sub>5</sub>, Sterimol L, buried Sterimol B<sub>5</sub>, and buried Sterimol L. These descriptors quantify the steric size of a substituent in Å, and are commonly employed for Quantitative Structure-Activity Relationship (QSAR) modeling. The buried Sterimol variants describe the steric effects within the first coordination sphere of a metal [57].

**Objectives.** As in the Drugs-75K tasks, the goal is to predict the Boltzmann-averaged value of each property across the conformer ensemble from the molecular graph  $G$ , a single conformer  $C_i \in \mathcal{C}$ , or the set  $\mathcal{C}$ .

**Dataset preparation.** In this study, we utilize the original 3D geometry structures of molecules and their corresponding Boltzmann-averaged properties provided in the Kraken dataset. Among the 78 physical-organic properties listed in the original dataset, we select four properties that demonstrate high variance across conformer ensembles, as illustrated in Figure S1.

**Data availability and license.** The Kraken dataset is publicly accessible at <https://kraken.cs.toronto.edu>. Its copyright is retained by the original authors. Under the permission of the original authors, the Kraken dataset with the conformer ensembles and the four conformer-level descriptors used in this study can be accessed at <https://github.com/SXKDZ/MARCEL/tree/main/dataset/Kraken>.

## A.3 EE

EE [34] is a dataset of 872 catalyst-substrate pairs involving 253 Rhodium (Rh)-bound atropisomeric catalysts derived from chiral bisphosphine, with 10 enamides as substrates. The dataset includes conformations of catalyst-substrate transition state complexes in two separate pro-S and pro-R configurations. The task is to predict the Enantiomeric Excess (EE) of the chemical reaction involving the substrate, defined as the absolute ratio between the concentration of each enantiomer in the product distribution.

**Objectives.** EE depends on the conformer ensembles of *each* pro-R and pro-S complex. The goal is to predict EE from the graphs of the catalyst and substrate ( $G_{\text{cat}}$ ,  $G_{\text{sub}}$ ), a conformer  $C_i^{(R)} \in \mathcal{C}^{(R)}$  and  $C_i^{(S)} \in \mathcal{C}^{(S)}$  for each complex, or the ensembles  $\mathcal{C}^{(R)}$  and  $\mathcal{C}^{(S)}$ .

**Dataset preparation.** The conformer ensembles are generated with Q2MM, which automatically generates Transition State Force Fields (TSFFs) in order to simulate the conformer ensembles of each

prochiral transition state complex. Then, the EE values are computed from the conformer ensembles by Boltzmann-averaging the activation energies for the competing transition states [34, 35]. Finally, we conduct the same conformer deduplication process as described in Section 3 with a RMSD threshold of 1.0.

**Data availability and license.** As of now, the EE dataset is proprietary, given that the publication addressing the conformer ensembles is still under preparation. Therefore, access to the EE dataset is restricted to review purposes only. We anticipate making the EE dataset publicly accessible following the acceptance of the corresponding paper.

#### A.4 BDE

BDE [36] is a dataset containing 5,915 organometallic catalysts  $ML_1L_2$  consisting of a metal center ( $M = Pd, Pt, Au, Ag, Cu, Ni$ ) coordinated to two flexible organic ligands ( $L_1$  and  $L_2$ ), each selected from a 91-membered ligand library. The data includes conformations of each unbound catalyst, as well as conformations of the catalyst when bound to ethylene and bromide after oxidative addition with vinyl bromide. Each catalyst has an electronic binding energy, computed as the difference in the minimum energies of the bound-catalyst complex and unbound catalyst, following the DFT-optimization of their respective conformer ensembles.

Although the binding energies are computed via DFT, the conformers provided for modeling are initially generated with Open Babel [37], followed by further geometric optimization steps, which ensures that the generated 3D structures are likely to be the global minimum energy conformers at the force field level [36, Supplementary Information]. We also note that obtaining DFT-optimized conformers for BDE is not feasible given the time-consuming nature of the process — a single geometric search using DFT can take 2 to 3 days. Therefore, this realistically represents the setting in which precise conformer ensembles are unknown at inference.

**Objectives.** The task is to predict the binding energy from the graphs of the unbound and bound catalyst, sampled conformers  $C_i^{(unbound)} \in \mathcal{C}^{(unbound)}$  and  $C_i^{(bound)} \in \mathcal{C}^{(bound)}$ , or the ensembles  $\mathcal{C}^{(unbound)}$  and  $\mathcal{C}^{(bound)}$ .

**Dataset preparation.** We employ Open Babel [37] to produce conformers for each unbound catalyst and each bound complex. In order to avoid redundancy, we follow a deduplication process as outlined in Section 3. For the unbound catalysts, a RMSD threshold value of 0.5 is applied, whereas for the bound complexes, a threshold of 1.0 is used.

**Data availability and license.** The binding energy descriptors can be accessed at <https://archive.materialscloud.org/record/2018.0014/v1> under the Creative Commons Attribution 4.0 International license. The conformers are publicly available at <https://github.com/SXKDZ/MARCEL/tree/main/datasets/BDE> under the Apache license.

## B IMPLEMENTATION DETAILS

### B.1 IMPLEMENTATION OF 1D MODELS

For the random forest model that operates on fingerprints, we employ three molecular fingerprint schemes: the Molecular ACCess System (MACCS) [42], Extended-Connectivity Fingerprints (ECFP) [41], and RDKit topological fingerprints [38]. Then, we concatenate their outputs into a single vector, which might lead to some feature redundancy, given the possible overlaps in these three fingerprint representations of the molecular structure. To tackle this issue, we remove any features that exhibit a high correlation exceeding 90% with the other features. For implementation, we employ Scikit-Learn [58] and compute fingerprints with RDKit [38].

For both LSTM and Transformer models that operate on SMILES strings, we use a Byte-Pair Encoding (BPE)-based tokenizer [59] that is pretrained on PubChem10M, which strikes a balance among character- and word-level representations and allows to handle large vocabularies in molecular corpora. For the Transformer model, we further follow the positional embedding scheme [44] to capture the positional relationship among tokens in the SMILES string.

Table S1: A summary of node and edge features used in 2D GNN models.

|      | Feature             | Explanation                                                                      |
|------|---------------------|----------------------------------------------------------------------------------|
| Node | AtomicNum           | Atomic number, representing the type of atom.                                    |
|      | ChiralTag           | Indicator of chirality, a property of asymmetry.                                 |
|      | TotalDegree         | Sum of implicit and explicit bonds of an atom.                                   |
|      | FormalCharge        | Charge of an atom assuming equal sharing of bonding electrons.                   |
|      | TotalNumHs          | Total number of hydrogen atoms bonded to the atom.                               |
|      | NumRadicalElectrons | Count of unpaired electrons in an atom.                                          |
|      | Hybridization       | Type of atomic orbital hybridization in the atom.                                |
|      | IsAromatic          | Boolean indicating if the atom is part of an aromatic ring.                      |
| Edge | IsInRing            | Boolean indicating if the atom is part of any ring structure.                    |
|      | BondType            | Type of the bond (e.g., single, double, triple, aromatic).                       |
|      | Stereo              | Stereochemistry of the bond (e.g., "none", "any", "Z", or "E" for double bonds). |
|      | IsConjugated        | Boolean indicating if the bond is part of a conjugated system.                   |

## B.2 FEATURIZATIONS OF MOLECULES FOR 2D MODELS

Following OGB [46], we employ a rich set of features for atoms (nodes) and bonds (edges) for 2D GNN models. A complete list of node and features can be found in Table S1.

## B.3 HYPERPARAMETER SPECIFICATIONS AND EXPERIMENTAL ENVIRONMENTS

Each model is trained over 2,000 epochs using the Adam optimizer [55] with early stopping triggered if there is no improvement in the training loss over 200 epochs. To ensure a fair comparison, the hidden dimension size is uniformly set to 128 for all models. Other hyperparameters mostly follow the original configurations as described in the respective papers. The complete hyperparameter set of each model can be found in [https://github.com/SXKDZ/MARCEL/tree/main/benchmarks/p\\_arams](https://github.com/SXKDZ/MARCEL/tree/main/benchmarks/p_arams).

We utilize PyTorch [60] and PyTorch-Geometric [61] to implement all deep learning models. Most of the experiments are conducted on servers equipped with Nvidia A100 GPUs, each with 40GB of memory. For memory-intensive models such as GemNet and LEFTNet, we use servers with Nvidia H100 GPUs, each with 80GB memory. The cumulative computation time across all experiments amounts to approximately 6,000 single GPU hours.

## C ADDITIONAL EXPERIMENTS ON EVALUATION SCHEMES OF THE CONFORMER SAMPLING STRATEGY

In this section, we conduct one additional experiment on the conformer ensemble learning strategies. We assess all 3D models on five tasks: Ionization Potential (IP) from the Drugs-75K dataset, B<sub>5</sub> and BurB<sub>5</sub> from the Kraken dataset, and tasks from the EE and BDE datasets.

In our previous setup, we evaluate the conformer sampling strategy using the lowest-energy conformer of each molecule at evaluation time, to provide a direct comparison to the single-conformer 3D models that are trained and tested with the lowest energy conformation. In these experiments, we continue to sample a random conformer uniformly from the conformer ensemble during training time, but consider two additional evaluation schemes: (1) evaluating model performance when encoding a randomly sampled conformer, and (2) evaluating model performance when averaging the per-conformer predictions across the entire conformer ensemble.

The results of these experiments are summarized in Table S2. In the table, we refer to the original evaluation scheme as "fixed", and the additional schemes as "random" and "all", respectively. We find that across all three schemes, using the lowest-energy conformer for evaluation consistently yields the best performance. This is expected, as the lowest-energy conformer contributes the most to ensemble-level descriptors. The random conformer evaluation scheme generally yields the worst performance, which is likely due to the introduction of noise from less relevant conformers at test

Table S2: Performance comparison of three conformer sampling variants with different evaluation strategies. All models are trained with a randomly sampled conformer from the ensemble. The last column summarizes the average rank across all datasets for each base model.

| Model     | Evaluation Strategy | Drugs-75K | Kraken         |                   | EE      | BDE    | Average Rank |
|-----------|---------------------|-----------|----------------|-------------------|---------|--------|--------------|
|           |                     | IP        | B <sub>5</sub> | BurB <sub>5</sub> |         |        |              |
| SchNet    | Fixed               | 0.4452    | 0.3235         | 0.2086            | 20.3595 | 1.9737 | 1            |
|           | Random              | 0.4498    | 0.3682         | 0.2454            | 22.0380 | 2.4416 | 3            |
|           | All                 | 0.4428    | 0.3856         | 0.2407            | 18.0296 | 2.0106 | 2            |
| DimeNet++ | Fixed               | 0.4395    | 0.3323         | 0.2237            | 15.0596 | 1.4741 | = 2          |
|           | Random              | 0.4555    | 0.3549         | 0.2222            | 13.5681 | 1.4688 | = 2          |
|           | All                 | 0.4479    | 0.3282         | 0.2001            | 12.3562 | 1.6270 | 1            |
| GemNet    | Fixed               | 0.4066    | 0.2694         | 0.1796            | 12.0541 | 1.6059 | 1            |
|           | Random              | 0.4250    | 0.4034         | 0.2534            | 16.1709 | 1.7894 | 3            |
|           | All                 | 0.4320    | 0.4523         | 0.2481            | 14.3952 | 1.6660 | 2            |
| PaiNN     | Fixed               | 0.4466    | 0.3441         | 0.2476            | 19.1521 | 1.9262 | 1            |
|           | Random              | 0.4770    | 0.3756         | 0.2478            | 21.3553 | 1.9411 | 3            |
|           | All                 | 0.4478    | 0.3458         | 0.2342            | 19.1955 | 1.8696 | 2            |
| ClofNet   | Fixed               | 0.4430    | 0.4524         | 0.2442            | 31.3733 | 2.5126 | 1            |
|           | Random              | 0.4530    | 0.4689         | 0.2736            | 31.3675 | 2.6310 | = 2          |
|           | All                 | 0.4363    | 0.4749         | 0.2855            | 34.3203 | 2.0271 | = 2          |
| LEFTNet   | Fixed               | 0.4149    | 0.2834         | 0.2120            | 20.3358 | 1.5276 | 1            |
|           | Random              | 0.4518    | 0.3177         | 0.2344            | 20.3740 | 1.5842 | 3            |
|           | All                 | 0.4274    | 0.3152         | 0.2170            | 18.8945 | 1.8663 | 2            |

time. Interestingly, we observe occasional performance improvement when averaging the predictions across all conformers in the ensemble, indicating that explicitly using ensemble-level information during evaluation can be beneficial.

## D RAW DATA

The raw performance data with standard deviation of Table 2 and Figure 3 is summarized in Table S3.

Table S3: Raw performance data (mean  $\pm$  standard deviation) of representative 1D, 2D, 3D, and conformer ensemble MRL models in terms of absolute test error.

| Category | Model                     | Drugs-75K                 |                           |                           | Kraken                    |                           |                           | EE                         | BDE                        |                            |                           |
|----------|---------------------------|---------------------------|---------------------------|---------------------------|---------------------------|---------------------------|---------------------------|----------------------------|----------------------------|----------------------------|---------------------------|
|          |                           | IP                        | EA                        | $\chi$                    | B <sub>5</sub>            | L                         | BurB <sub>5</sub>         |                            |                            | BurL                       |                           |
| 1D       | Random forest             | 0.4987 <sub>±0.0037</sub> | 0.4747 <sub>±0.0022</sub> | 0.2732 <sub>±0.0031</sub> | 0.4760 <sub>±0.0041</sub> | 0.4303 <sub>±0.0090</sub> | 0.2758 <sub>±0.0180</sub> | 0.1521 <sub>±0.0149</sub>  | 61.2963 <sub>±2.8640</sub> | 3.0335 <sub>±0.2693</sub>  |                           |
|          | LSTM                      | 0.4788 <sub>±0.0024</sub> | 0.4648 <sub>±0.0002</sub> | 0.2505 <sub>±0.0050</sub> | 0.4879 <sub>±0.0280</sub> | 0.5142 <sub>±0.0411</sub> | 0.2813 <sub>±0.0041</sub> | 0.1924 <sub>±0.0028</sub>  | 64.0088 <sub>±2.3708</sub> | 2.8279 <sub>±0.0728</sub>  |                           |
|          | Transformer               | 0.6617 <sub>±0.0023</sub> | 0.5850 <sub>±0.0031</sub> | 0.4073 <sub>±0.0006</sub> | 0.9611 <sub>±0.0813</sub> | 0.8389 <sub>±0.0431</sub> | 0.4929 <sub>±0.0369</sub> | 0.2781 <sub>±0.0207</sub>  | 62.0816 <sub>±2.1789</sub> | 10.0771 <sub>±0.6457</sub> |                           |
| 2D       | GIN                       | 0.4354 <sub>±0.0029</sub> | 0.4169 <sub>±0.0032</sub> | 0.2260 <sub>±0.0017</sub> | 0.3128 <sub>±0.0264</sub> | 0.4003 <sub>±0.0341</sub> | 0.1719 <sub>±0.0031</sub> | 0.1200 <sub>±0.0040</sub>  | 62.3065 <sub>±2.9010</sub> | 2.6368 <sub>±0.2276</sub>  |                           |
|          | GIN-VN                    | 0.4361 <sub>±0.0059</sub> | 0.4169 <sub>±0.0083</sub> | 0.2267 <sub>±0.0002</sub> | 0.3567 <sub>±0.0031</sub> | 0.4344 <sub>±0.0416</sub> | 0.2422 <sub>±0.0033</sub> | 0.1741 <sub>±0.0109</sub>  | 62.3815 <sub>±2.1882</sub> | 2.7417 <sub>±0.2446</sub>  |                           |
|          | ChemProp                  | 0.4595 <sub>±0.0028</sub> | 0.4417 <sub>±0.0045</sub> | 0.2441 <sub>±0.0012</sub> | 0.4850 <sub>±0.0068</sub> | 0.5452 <sub>±0.0454</sub> | 0.3002 <sub>±0.0086</sub> | 0.1948 <sub>±0.0138</sub>  | 61.0336 <sub>±2.9715</sub> | 2.6616 <sub>±0.1429</sub>  |                           |
| 3D       | GraphGPS                  | 0.4351 <sub>±0.0049</sub> | 0.4085 <sub>±0.0055</sub> | 0.2212 <sub>±0.0054</sub> | 0.3450 <sub>±0.0324</sub> | 0.4363 <sub>±0.0133</sub> | 0.2066 <sub>±0.0115</sub> | 0.1500 <sub>±0.0138</sub>  | 61.6251 <sub>±1.3743</sub> | 2.4827 <sub>±0.1992</sub>  |                           |
|          | SchNet                    | 0.4394 <sub>±0.0062</sub> | 0.4207 <sub>±0.0021</sub> | 0.2243 <sub>±0.0089</sub> | 0.3293 <sub>±0.0068</sub> | 0.5458 <sub>±0.0341</sub> | 0.2295 <sub>±0.0111</sub> | 0.1861 <sub>±0.0095</sub>  | 17.7421 <sub>±1.0899</sub> | 2.5488 <sub>±0.0050</sub>  |                           |
|          | DimeNet++                 | 0.4441 <sub>±0.0087</sub> | 0.4233 <sub>±0.0072</sub> | 0.2436 <sub>±0.0075</sub> | 0.3510 <sub>±0.0107</sub> | 0.4174 <sub>±0.0397</sub> | 0.2097 <sub>±0.0160</sub> | 0.1526 <sub>±0.0072</sub>  | 14.6414 <sub>±2.2791</sub> | 1.4503 <sub>±0.0370</sub>  |                           |
|          | GemNet                    | 0.4069 <sub>±0.0007</sub> | 0.3922 <sub>±0.0024</sub> | 0.1970 <sub>±0.0039</sub> | 0.2789 <sub>±0.0125</sub> | 0.3754 <sub>±0.0086</sub> | 0.1782 <sub>±0.0099</sub> | 0.1635 <sub>±0.0063</sub>  | 18.0338 <sub>±2.4777</sub> | 1.6530 <sub>±0.3081</sub>  |                           |
|          | PaiNN                     | 0.4505 <sub>±0.0041</sub> | 0.4495 <sub>±0.0054</sub> | 0.2324 <sub>±0.0040</sub> | 0.3443 <sub>±0.0388</sub> | 0.4471 <sub>±0.0324</sub> | 0.2395 <sub>±0.0176</sub> | 0.1673 <sub>±0.0088</sub>  | 20.2359 <sub>±1.2128</sub> | 2.1261 <sub>±0.0920</sub>  |                           |
|          | ClofNet                   | 0.4393 <sub>±0.0084</sub> | 0.4251 <sub>±0.0066</sub> | 0.2378 <sub>±0.0020</sub> | 0.4873 <sub>±0.0093</sub> | 0.6417 <sub>±0.0362</sub> | 0.2884 <sub>±0.0166</sub> | 0.2529 <sub>±0.0052</sub>  | 33.9473 <sub>±1.4633</sub> | 2.6057 <sub>±0.0236</sub>  |                           |
|          | LEFTNet                   | 0.4174 <sub>±0.0007</sub> | 0.3964 <sub>±0.0009</sub> | 0.2083 <sub>±0.0054</sub> | 0.3072 <sub>±0.0012</sub> | 0.4493 <sub>±0.0261</sub> | 0.2176 <sub>±0.0010</sub> | 0.1486 <sub>±0.0095</sub>  | 19.7974 <sub>±1.4097</sub> | 1.5328 <sub>±0.0567</sub>  |                           |
|          | SchNet                    | 0.4452 <sub>±0.0080</sub> | 0.4232 <sub>±0.0042</sub> | 0.2243 <sub>±0.0022</sub> | 0.3235 <sub>±0.0147</sub> | 0.4598 <sub>±0.0041</sub> | 0.2086 <sub>±0.0111</sub> | 0.1739 <sub>±0.0142</sub>  | 20.3595 <sub>±1.5260</sub> | 1.9737 <sub>±0.0125</sub>  |                           |
|          | DimeNet++                 | 0.4395 <sub>±0.0032</sub> | 0.4217 <sub>±0.0040</sub> | 0.2432 <sub>±0.0048</sub> | 0.3323 <sub>±0.0320</sub> | 0.4153 <sub>±0.0208</sub> | 0.2237 <sub>±0.0122</sub> | 0.1561 <sub>±0.0241</sub>  | 15.0596 <sub>±0.2867</sub> | 1.4741 <sub>±0.0349</sub>  |                           |
|          | GemNet                    | 0.4066 <sub>±0.0015</sub> | 0.3910 <sub>±0.0004</sub> | 0.2027 <sub>±0.0013</sub> | 0.2694 <sub>±0.0221</sub> | 0.3488 <sub>±0.0252</sub> | 0.1796 <sub>±0.0098</sub> | 0.1184 <sub>±0.0033</sub>  | 12.0541 <sub>±0.7735</sub> | 1.6059 <sub>±0.1094</sub>  |                           |
|          | PaiNN                     | 0.4466 <sub>±0.0087</sub> | 0.4393 <sub>±0.0045</sub> | 0.2331 <sub>±0.0037</sub> | 0.3441 <sub>±0.0161</sub> | 0.4358 <sub>±0.0343</sub> | 0.2476 <sub>±0.0070</sub> | 0.1543 <sub>±0.0022</sub>  | 19.1521 <sub>±0.2386</sub> | 1.9262 <sub>±0.0188</sub>  |                           |
|          | ClofNet                   | 0.4430 <sub>±0.0074</sub> | 0.4237 <sub>±0.0005</sub> | 0.2335 <sub>±0.0090</sub> | 0.4524 <sub>±0.0935</sub> | 0.5962 <sub>±0.0074</sub> | 0.2442 <sub>±0.0109</sub> | 0.1756 <sub>±0.0112</sub>  | 31.3733 <sub>±1.9892</sub> | 2.5126 <sub>±0.2366</sub>  |                           |
| LEFTNet  | 0.4149 <sub>±0.0019</sub> | 0.3988 <sub>±0.0048</sub> | 0.2141 <sub>±0.0084</sub> | 0.2834 <sub>±0.0068</sub> | 0.4407 <sub>±0.0531</sub> | 0.2120 <sub>±0.0097</sub> | 0.1547 <sub>±0.0101</sub> | 20.3358 <sub>±0.6614</sub> | 1.5276 <sub>±0.0088</sub>  |                            |                           |
| Ensemble | SchNet                    | Mean                      | 0.4583 <sub>±0.0019</sub> | 0.4410 <sub>±0.0018</sub> | 0.2371 <sub>±0.0098</sub> | 0.3075 <sub>±0.0151</sub> | 0.4691 <sub>±0.0234</sub> | 0.2282 <sub>±0.0206</sub>  | 0.1619 <sub>±0.0062</sub>  | 20.1392 <sub>±1.5748</sub> | 2.5312 <sub>±0.0248</sub> |
|          |                           | DeepSet                   | 0.4537 <sub>±0.0065</sub> | 0.4396 <sub>±0.0010</sub> | 0.2385 <sub>±0.0066</sub> | 0.3105 <sub>±0.0381</sub> | 0.4322 <sub>±0.0464</sub> | 0.2249 <sub>±0.0234</sub>  | 0.1535 <sub>±0.0076</sub>  | 18.0495 <sub>±1.2846</sub> | 2.2941 <sub>±0.2229</sub> |
|          |                           | Attention                 | 0.4556 <sub>±0.0075</sub> | 0.4382 <sub>±0.0125</sub> | 0.2380 <sub>±0.0007</sub> | 0.2704 <sub>±0.0187</sub> | 0.4517 <sub>±0.0132</sub> | 0.2024 <sub>±0.0183</sub>  | 0.1443 <sub>±0.0043</sub>  | 14.2238 <sub>±0.5451</sub> | 2.6445 <sub>±0.0031</sub> |
|          | DimeNet++                 | Mean                      | 0.4488 <sub>±0.0086</sub> | 0.4340 <sub>±0.0079</sub> | 0.2425 <sub>±0.0060</sub> | 0.2630 <sub>±0.0122</sub> | 0.3828 <sub>±0.0331</sub> | 0.1960 <sub>±0.0059</sub>  | 0.1268 <sub>±0.0060</sub>  | 12.0259 <sub>±0.8933</sub> | 1.7964 <sub>±0.1260</sub> |
|          |                           | DeepSet                   | 0.4126 <sub>±0.0076</sub> | 0.3944 <sub>±0.0034</sub> | 0.2267 <sub>±0.0047</sub> | 0.2889 <sub>±0.0069</sub> | 0.3468 <sub>±0.0090</sub> | 0.1783 <sub>±0.0110</sub>  | 0.1339 <sub>±0.0087</sub>  | 15.5754 <sub>±2.6294</sub> | 1.7533 <sub>±0.0163</sub> |
|          |                           | Attention                 | 0.4188 <sub>±0.0024</sub> | 0.4030 <sub>±0.0075</sub> | 0.2325 <sub>±0.0028</sub> | 0.3718 <sub>±0.0300</sub> | 0.3628 <sub>±0.0259</sub> | 0.1899 <sub>±0.0081</sub>  | 0.1185 <sub>±0.0105</sub>  | 13.3643 <sub>±1.4309</sub> | 2.5714 <sub>±0.2149</sub> |
|          | GemNet                    | Mean                      | 0.4505 <sub>±0.0052</sub> | 0.4334 <sub>±0.0023</sub> | 0.2289 <sub>±0.0032</sub> | 0.2635 <sub>±0.0053</sub> | 0.3753 <sub>±0.0036</sub> | 0.1671 <sub>±0.0154</sub>  | 0.1587 <sub>±0.0029</sub>  | 11.6142 <sub>±1.7271</sub> | 2.1914 <sub>±0.0605</sub> |
|          |                           | DeepSet                   | 0.4187 <sub>±0.0022</sub> | 0.4002 <sub>±0.0012</sub> | 0.2169 <sub>±0.0036</sub> | 0.2313 <sub>±0.0026</sub> | 0.3386 <sub>±0.0269</sub> | 0.1589 <sub>±0.0068</sub>  | 0.0947 <sub>±0.0012</sub>  | 13.9273 <sub>±1.8656</sub> | 2.2532 <sub>±0.2106</sub> |
|          |                           | Attention                 | 0.4212 <sub>±0.0017</sub> | 0.4221 <sub>±0.0097</sub> | 0.2260 <sub>±0.0056</sub> | 0.2670 <sub>±0.0026</sub> | 0.3554 <sub>±0.0147</sub> | 0.1769 <sub>±0.0153</sub>  | 0.1346 <sub>±0.0075</sub>  | 12.0249 <sub>±1.8418</sub> | 2.6810 <sub>±0.0223</sub> |
|          | PaiNN                     | Mean                      | 0.4591 <sub>±0.0024</sub> | 0.4425 <sub>±0.0064</sub> | 0.2360 <sub>±0.0032</sub> | 0.2877 <sub>±0.0252</sub> | 0.3950 <sub>±0.0233</sub> | 0.1817 <sub>±0.0091</sub>  | 0.1472 <sub>±0.0039</sub>  | 16.4239 <sub>±0.0743</sub> | 1.8744 <sub>±0.1657</sub> |
|          |                           | DeepSet                   | 0.4471 <sub>±0.0071</sub> | 0.4269 <sub>±0.0033</sub> | 0.2294 <sub>±0.0065</sub> | 0.2225 <sub>±0.0218</sub> | 0.3619 <sub>±0.0192</sub> | 0.1693 <sub>±0.0111</sub>  | 0.1324 <sub>±0.0091</sub>  | 13.5570 <sub>±0.5905</sub> | 2.2097 <sub>±0.0586</sub> |
|          |                           | Attention                 | 0.4641 <sub>±0.0016</sub> | 0.4567 <sub>±0.0094</sub> | 0.2471 <sub>±0.0049</sub> | 0.3496 <sub>±0.0140</sub> | 0.4109 <sub>±0.0167</sub> | 0.2123 <sub>±0.0005</sub>  | 0.1506 <sub>±0.0029</sub>  | 19.1556 <sub>±2.2765</sub> | 2.2335 <sub>±0.1255</sub> |
|          | ClofNet                   | Mean                      | 0.4536 <sub>±0.0030</sub> | 0.4301 <sub>±0.0007</sub> | 0.2365 <sub>±0.0075</sub> | 0.3555 <sub>±0.0193</sub> | 0.4485 <sub>±0.0053</sub> | 0.2473 <sub>±0.0076</sub>  | 0.2022 <sub>±0.0212</sub>  | 19.9710 <sub>±0.7745</sub> | 2.0106 <sub>±0.0856</sub> |
|          |                           | DeepSet                   | 0.4280 <sub>±0.0056</sub> | 0.4033 <sub>±0.0024</sub> | 0.2199 <sub>±0.0073</sub> | 0.3228 <sub>±0.0020</sub> | 0.4742 <sub>±0.0161</sub> | 0.2263 <sub>±0.0249</sub>  | 0.1548 <sub>±0.0039</sub>  | 13.9647 <sub>±1.2753</sub> | 2.3576 <sub>±0.0496</sub> |
|          |                           | Attention                 | 0.4330 <sub>±0.0071</sub> | 0.4107 <sub>±0.0048</sub> | 0.2220 <sub>±0.0084</sub> | 0.3734 <sub>±0.0267</sub> | 0.4963 <sub>±0.0286</sub> | 0.2178 <sub>±0.0186</sub>  | 0.1690 <sub>±0.0281</sub>  | 26.7133 <sub>±1.7225</sub> | 2.6652 <sub>±0.1438</sub> |
|          | LEFTNet                   | Mean                      | 0.4402 <sub>±0.0062</sub> | 0.4267 <sub>±0.0026</sub> | 0.2183 <sub>±0.0007</sub> | 0.2949 <sub>±0.0001</sub> | 0.3643 <sub>±0.0352</sub> | 0.2098 <sub>±0.0146</sub>  | 0.1386 <sub>±0.0007</sub>  | 18.9245 <sub>±2.0136</sub> | 2.0440 <sub>±0.0076</sub> |
|          |                           | DeepSet                   | 0.4167 <sub>±0.0043</sub> | 0.3953 <sub>±0.0000</sub> | 0.2069 <sub>±0.0022</sub> | 0.2644 <sub>±0.0130</sub> | 0.3866 <sub>±0.0270</sub> | 0.2023 <sub>±0.0026</sub>  | 0.1441 <sub>±0.0042</sub>  | 18.4189 <sub>±1.8922</sub> | 2.5165 <sub>±0.3077</sub> |
|          |                           | Attention                 | 0.4229 <sub>±0.0059</sub> | 0.4067 <sub>±0.0047</sub> | 0.2198 <sub>±0.0011</sub> | 0.3161 <sub>±0.0116</sub> | 0.4324 <sub>±0.0292</sub> | 0.2017 <sub>±0.0023</sub>  | 0.1508 <sub>±0.0075</sub>  | 18.9988 <sub>±1.6904</sub> | 2.6361 <sub>±0.1560</sub> |