

REFINE DRUGS, DON’T COMPLETE THEM: UNIFORM-SOURCE DISCRETE FLOWS FOR FRAGMENT-BASED DRUG DISCOVERY

Benno Kaech*

Luis Wyss[†]

Karsten Borgwardt[†]

Gianvito Grasso*

ABSTRACT

We introduce InVirtuoGen, a discrete flow generative model for fragmented SMILES for *de novo* and fragment-constrained generation, and target-property/lead optimization of small molecules. The model learns to transform a uniform source over all possible tokens into the data distribution. Unlike masked models, its training loss accounts for predictions on all sequence positions at every denoising step, shifting the generation paradigm from completion to refinement, and decoupling the number of sampling steps from the sequence length. For *de novo* generation, InVirtuoGen achieves a stronger quality-diversity pareto frontier than prior fragment-based models and competitive performance on fragment-constrained tasks. For property and lead optimization, we propose a hybrid scheme that combines a genetic algorithm with a Proximal Property Optimization fine-tuning strategy adapted to discrete flows. Our approach sets a new state-of-the-art on the Practical Molecular Optimization benchmark, measured by top-10 AUC across tasks, and yields higher docking scores in lead optimization than previous baselines. InVirtuoGen thus establishes a versatile generative foundation for drug discovery, from early hit finding to multi-objective lead optimization. We further contribute to open science by releasing pretrained checkpoints and code, making our results fully reproducible¹.

1 INTRODUCTION

Fragment-based drug discovery (FBDD) is widely used in both academia and industry for its efficient exploration of chemical space (Kirkman et al., 2024). FBDD relies on fragment-constrained design, in which new candidate molecules are generated by preserving specific substructures, such as active scaffolds or pharmacophores, and modifying surrounding regions to tune properties (Kirsch et al., 2019). However, FBDD is typically guided by expert-defined heuristics based on chemical intuition, limiting exploration of the vast chemical space. In contrast, *in silico* molecular generation seeks to formalize and automate this intuition, leveraging data-driven generative models to navigate chemical space more systematically. While many such models have emerged as promising candidates to accelerate the drug discovery pipeline (Zeng et al., 2022), their adoption in practice remains limited. A key barrier is that their molecular representations are often poorly aligned with medicinal chemistry workflows, making them difficult to integrate into existing pipelines. Although graphs provide a natural representation for molecules, generative frameworks tailored to graph-structured data remain limited in performance. For instance, the state-of-the-art AutoGraph (Chen et al., 2025) sidesteps direct graph generation by linearizing graphs into sequences via depth-first traversal and relying on next token prediction. Similarly, the Simplified Molecular Input Line Entry System (SMILES) (Weininger, 1988) is commonly used when working with small molecules. SMILES encodes molecular graphs as sequences via depth-first traversal of a spanning tree with annotations for branches and ring closures. However, both linearizations disrupt chemically meaningful substructures, offering limited control over scaffold retention or fragment assembly and making them poorly suited for fragment-based drug discovery (Jinsong et al., 2024). We propose InVirtuoGen, a continuous-time discrete flow model (Campbell et al., 2024; Gat et al., 2024) that operates directly on fragmented SMILES.

*In Virtuo Laboratories. Correspondence: info@invirtuolabs.com

[†]Max Planck Institute of Biochemistry

¹https://github.com/invirtuolabs/InVirtuoGen_results

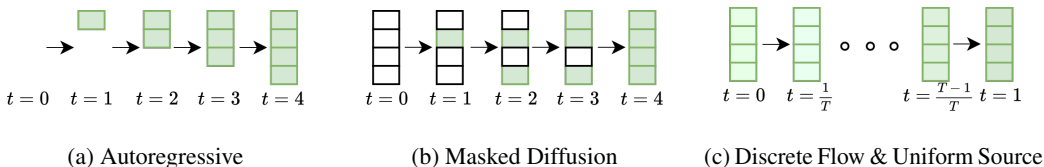


Figure 1: Comparison of generation paradigms: (a) autoregressive models generate tokens sequentially (here simplified by omitting BOS/EOS tokens), (b) masked diffusion models iteratively reveal masked positions, and (c) discrete flows refine all positions starting from a uniform source distribution, where shading indicates the transition from random tokens to data.

2 RELATED WORK

While numerous generative models have been proposed for small-molecule drug design (Jensen, 2019; Olivecrona et al., 2017; Morrison et al., 2024; Gao et al., 2022b; Nigam et al., 2020; Bou et al., 2024), few are explicitly designed for fragment-level control. In this work, we focus on approaches that operate on sequential fragment-based representations. Several alternative approaches construct molecules via graph-based operations on substructures, for example by adding or deleting fragments through Markov Chain Monte Carlo sampling (Xie et al., 2021), using graph-based Variational Auto-Encoders conditioned on identified substructures (Jin et al., 2020; 2018; Maziarz et al., 2024), or applying graph-based genetic algorithms (GA) (Jensen, 2019; Tripp and Hernández-Lobato, 2023). Although these models encode domain-specific priors, they often suffer from poor scalability and limited generalization beyond known chemical regions, in part due to their reliance on graph operations and discrete mutation strategies. By contrast, generative models operating on linear sequential representations of fragmented SMILES, such as SAFE-GPT (Noutahi et al., 2023) and GenMol (Lee et al., 2025), offer a more scalable and expressive alternative, and form the primary baselines for our work².

Autoregressive Models Autoregressive approaches, such as SAFE-GPT (Noutahi et al., 2023), generate molecular sequences token by token in a fixed left-to-right order. This ordering is arbitrary with respect to molecular structure, which is inherently unordered³.

Masked Diffusion Models Masked discrete diffusion models, such as GenMol (Lee et al., 2025), iteratively unmask tokens starting from a fully masked input. While predictions are produced for the entire sequence at each denoising step, the training objective accounts for errors only on the masked positions. As a consequence, once a token is unmasked during sampling, it is treated as fixed and no longer updated. This introduces a fundamental limitation: the number of sampling steps is bounded by the number of initially masked tokens, unless arbitrary remasking heuristics are included.

Our Contributions Our method departs from completion-style generation and instead refines all positions simultaneously at every denoising step (Figure 1). This training paradigm enables coordinated updates across the molecule and decouples sampling steps from sequence length, aligning with our central principle: *refine drugs, don’t complete them*. Concurrent to our work, Schiff et al. (Schiff et al., 2025) proposed Uniform Discrete Language Models, which similarly allow simultaneous token updates but remain within a diffusion-based framework. Concretely, we present the first discrete flow model for fragmented SMILES with a refinement-based training paradigm, show state-of-the-art performance in *de novo* generation and competitive performance on fragment-constrained generation tasks, and introduce a hybrid optimization framework combining Proximal Property Optimization (Schulman et al., 2017), adapted to discrete flows, with a genetic algorithm. Our framework achieves state-of-the-art results on the Practical Molecular Optimization (PMO) benchmark (Gao et al., 2022a) and improved results in lead optimization over prior baselines (Lee et al., 2025).

²We compare against GenMol using the results reported in their paper. While the source code is publicly available, to the best of our knowledge, the pretrained checkpoints are only distributed through NVIDIA NIM/NGC under the NVIDIA Open Model License, which currently does not allow unrestricted download. As a result, we were unable to run GenMol directly and rely on the published numbers for comparison.

³Although, there exist a canonical SMILES encoding scheme, it is still an arbitrary imposed ordering.

2.1 DISCRETE FLOW MODELS

We adopt the discrete flow model framework of Gat et al. (2024), where the goal is to transform samples from a source distribution $X_0 \sim p$ into samples from a target distribution $X_1 \sim q$. Training data consists of interpolation pairs (X_0, X_1) , sampled independently from the source and target. We choose the linear scheduler κ_j^t of Gat et al. (2024), resulting in following the probability path:

$$p_t(x^i | x_0, x_1) = (1 - t) \delta_{x_0}(x^i) + t \delta_{x_1}(x^i), \quad t \in [0, 1]. \quad (1)$$

During sampling, each token X_t^i is updated according to the discrete-time Markov update

$$X_{t+h}^i \sim \delta_{X_t^i}(\cdot) + h u_t^i(\cdot, X_t), \quad (2)$$

where u_t is the *probability velocity* and $h > 0$ is the step size. Following Gat et al. (2024), u_t must satisfy the validity constraints

$$\sum_{x^i \in [d]} u_t^i(x^i, z) = 0, \quad u_t^i(x^i, z) \geq 0 \text{ for all } i \text{ and } x^i \neq z. \quad (3)$$

For our scheduler choice, the training objective becomes

$$\mathcal{L}(\theta) = -\mathbb{E}_{t \sim U(0,1), (X_0, X_1), X_t} \frac{1}{1-t^2} \sum_i \log p_{1|T}(X_1^i | X_t), \quad (4)$$

where $p_{1|T}$ denotes the model prediction and the sum is over the sequence. The time-dependent loss weighting term, not present in its original formulation, was inspired by Sahoo et al. (2024) and places greater emphasis on later timesteps, encouraging higher accuracy near the end of the trajectory. As a backbone model, we use a diffusion transformer (Peebles and Xie, 2022) to parameterize $p_{1|T}$, leveraging its bidirectional self-attention to capture long-range dependencies between fragments while predicting the target token distribution at each position. More details about our experimental setup is given in Appendix A.1.

2.2 FRAGMENTED SMILES NOTATION & PREPROCESSING

Our molecule representation is based on and marginally extends the Sequential Attachment-Based Fragment Embedding (SAFE) framework of Noutahi et al. (2023) by encoding molecules as sequences of fragment blocks with explicit attachment points, improving readability and direct control over molecular substructures. To produce chemically meaningful fragments, the molecules are decomposed using the revised BRICS algorithm (Degen et al., 2008), with the locations of bond breaks marked by attachment points of the form $[i*]$, where i enumerates the broken bonds and we separate fragments with spaces. In Fig. 2, we illustrate the difference between the SMILES, SAFE, and our notation. To remove any implicit ordering bias, the resulting fragments are randomly shuffled rather than ordered by their attachment point in the original molecule. The resulting fragmented SMILES strings are tokenized at the atomic level⁴, yielding a vocabulary of 204 tokens.

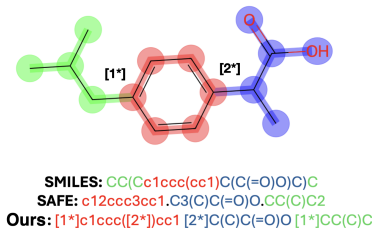


Figure 2: Comparison between SMILES, SAFE, and our notation for the same molecule. Our notation preserves fragment integrity while providing explicit attachment point numbering that facilitates bidirectional modeling of molecular structure.

3 EXPERIMENTS

We evaluate InVirtuoGen on four molecular design tasks: (i) *de novo* generation of diverse, synthesizable, drug-like molecules; (ii) fragment-constrained design with predefined scaffolds or pharmacophores; (iii) target property optimization on the PMO benchmark (Gao et al., 2022a), assessing

⁴Note that we do not tokenize e.g. ‘Cl’ as two tokens, but as a single one. The same holds for the attachment point tokens.

sample efficiency and oracle performance; and (iv) lead optimization, optimizing docking scores under similarity, synthesizability and drug-likeness constraints.

For a fair comparison, we pretrain on the same datasets as GenMol (Lee et al., 2025) and SAFE-GPT (El Mesbahi and Noutahi, 2024): ZINC (Sterling and Irwin, 2015) and UniChem (Chambers et al., 2013), containing roughly one billion molecules.

The non-autoregressive formulation enables bidirectional attention, making the model well-suited for the inherently unordered representation. However, this also implies that we explicitly decouple sequence length from the generation process by factorizing the output distribution as

$$p_{\theta}(\mathbf{x}) = p(n) p_{\theta}(\mathbf{x} \mid n), \quad (5)$$

where $p(n)$ models the sequence length. To compare with GenMol, our base distribution is also chosen to be the empirical length distribution of ZINC250k, a curated subset of ZINC (Sterling and Irwin, 2015) containing synthesizable, drug-like compounds.

3.1 De Novo GENERATION

For drug discovery, generated molecules must be diverse, synthesizable, and drug-like. We evaluate these aspects with four metrics, following prior work on fragmented SMILES: *Validity* (fraction of valid SMILES), *Uniqueness* (fraction of unique valid molecules), *Diversity* (average Tanimoto distance of Morgan fingerprints (Polykovskiy et al., 2018; Rogers and Hahn, 2010)), and *Quality* (Lee et al., 2025) (fraction of valid, unique molecules with $\text{QED} \geq 0.6$ (Bickerton et al., 2012) and $\text{SA} \leq 4$ (Ertl and Schuffenhauer, 2009)). Quality and diversity are the main criteria but trade off against each other. As in GenMol, we tune this balance via the softmax temperature T and a noise scale r (modulating Gumbel noise⁵). During generation, r is damped as $(1 - t)$ and T is annealed, promoting early exploration and late refinement. Empirically, sampling directly from the predicted token-wise distribution

$$X_{t+h}^i \sim \hat{p}_t^i(X_t), \quad (6)$$

outperforms Eq. 2 significantly. Importantly, our update departs from masked discrete diffusion models: instead of replacing only masked tokens, all sequence positions possibly change simultaneously at every step, leading to a refinement process rather than arbitrary order completion. We provide more details and investigate the benefits of our sampling method in Appendix B.1.

Results In Figure 3, we present and compare our results to other state-of-the-art fragment-based generative models. Increasing the time-granularity (i.e., using smaller timestep sizes h) consistently improves both quality and diversity. InVirtuoGen consistently achieves a superior pareto frontier, with the largest gains at high time-granularity ($h = 0.001$), outperforming all baselines across the full quality-diversity spectrum. As described in Appendix B.1, sampling with $h = 0.01$ yields a comparable number of model calls, while still showing a modest performance increase over GenMol, particularly in the lower-diversity regime. We provide the ZINC250k sequence length distribution, non-curated samples, timing studies and additional results in Appendix B.1, including the pareto frontier obtained by sampling according to Eq. 2.

3.2 FRAGMENT-CONSTRAINED GENERATION

A core task in FBDD is to generate molecules given one or more fragments the molecule should contain. We follow the benchmark of Noutahi et al. (2023), which uses fragments from ten known drugs, and evaluate five subtasks: linker design (connecting two or more terminal fragments with a feasible linker), scaffold morphing (modifying the core scaffold while preserving pharmacophoric features), motif extension (growing a fixed motif with new substituents), scaffold decoration (attaching functional groups at predefined positions), and superstructure generation (assembling multiple fragments into a coherent larger molecule). In our notation, as for GenMol, linker design and scaffold morphing yield identical prompts and thus identical results. Note that both next-token prediction models and masked discrete diffusion models such as GenMol can be prompted in a straightforward manner by using a prefix. In contrast, our model predicts a potential completely different sequence at every timestep. Therefore, to ensure that the generated molecules remain consistent with given fragments, the corresponding positions are naively overwritten at every timestep of the simulation. We note that fragment-constrained generation is fundamentally at odds with our refinement philosophy. It requires fixing certain positions, preventing the holistic refinement that makes our approach powerful. We include these results for comparison with prior work, but did not optimize for this setting as it essentially reduces our method to a masked completion approach, negating our core contribution.

⁵GenMol instead perturbs the order of the confidence-based token unmasking

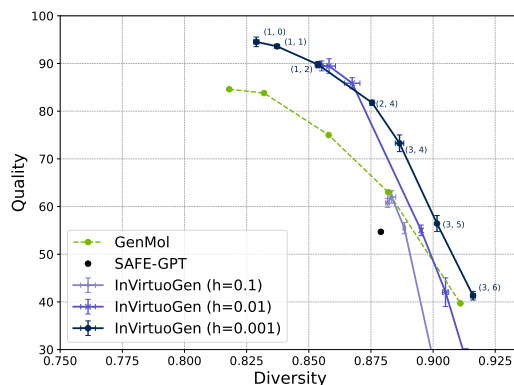


Figure 3: Quality-diversity trade-off for GenMol, SAFE-GPT (single point, as no quality-diversity scan data is available), and our model at different simulation time granularities ($h \in \{0.1, 0.01, 0.001\}$). Curves correspond to varying sampling noise (T, r) , where T is the softmax temperature and r is the Gumbel noise scale.

Results In Table 1, we report results for the fragment-constrained generation tasks. Averaging over all tasks and three random seeds, our model yields competitive results in terms of quality and diversity. Additionally, for GenMol, the generation parameters were tuned separately for each task, whereas we used a single fixed parameter setting across all tasks, more details in Appendix A.2. We provide non-curated samples for the different tasks in the Appendix B.2.

Table 1: Performance across five fragment-constrained generation tasks, averaged over three random seeds: motif extension, linker design, superstructure generation, scaffold morphing, and scaffold decoration. The results for GenMol (Lee et al., 2025) and SAFE-GPT (Noutahi et al., 2023) are taken from Lee et al. (2025).

Task	Method	Diversity	Quality	Uniqueness	Validity
Motif Extension	SAFE-GPT	0.56 ± 0.003	18.60 ± 2.100	66.80 ± 1.200	96.10 ± 1.900
	GenMol	0.62 ± 0.002	30.10 ± 0.400	77.50 ± 0.100	82.90 ± 0.100
	InVirtuoGen	0.62 ± 0.001	39.17 ± 0.665	96.17 ± 0.308	68.43 ± 1.212
Linker Design	SAFE-GPT	0.55 ± 0.007	21.70 ± 1.100	82.50 ± 1.900	76.60 ± 5.100
	GenMol	0.55 ± 0.002	21.90 ± 0.400	83.70 ± 0.500	100.00
	InVirtuoGen	0.51 ± 0.002	21.97 ± 1.053	83.89 ± 1.781	60.67 ± 1.642
Scaffold Morphing	SAFE-GPT	0.51 ± 0.011	16.70 ± 2.300	70.40 ± 5.700	58.90 ± 6.800
	GenMol	0.55 ± 0.002	21.90 ± 0.400	83.70 ± 0.500	100.00
	InVirtuoGen	0.51 ± 0.002	21.97 ± 1.053	83.89 ± 1.781	60.67 ± 1.642
Superstructure Design	SAFE-GPT	0.57 ± 0.028	14.30 ± 3.700	83.00 ± 5.900	95.70 ± 2.000
	GenMol	0.60 ± 0.009	34.80 ± 1.000	83.60 ± 1.000	97.50 ± 0.900
	InVirtuoGen	0.73 ± 0.001	28.63 ± 1.226	99.67 ± 0.055	74.90 ± 0.432
Scaffold Decoration	SAFE-GPT	0.57 ± 0.008	10.00 ± 1.400	74.70 ± 2.500	97.70 ± 0.300
	GenMol	0.59 ± 0.001	31.80 ± 0.500	82.70 ± 1.800	96.60 ± 0.800
	InVirtuoGen	0.56 ± 0.004	35.90 ± 0.942	88.21 ± 0.546	90.47 ± 0.754
Average	SAFE-GPT	0.55 ± 0.006	16.26 ± 1.031	75.48 ± 1.773	85.00 ± 1.788
	GenMol	0.58 ± 0.002	28.10 ± 0.263	82.24 ± 0.436	95.40 ± 0.242
	InVirtuoGen	0.59 ± 0.001	29.53 ± 0.600	90.37 ± 0.432	71.03 ± 0.539

3.3 TARGET PROPERTY OPTIMIZATION

The PMO benchmark evaluates molecular optimization under conditions that mirror practical drug discovery, with limited oracle calls and diverse pharmacologically relevant objectives. We introduce a hybrid approach that combines a genetic algorithm, which provides fast convergence through recombination of high-scoring molecules, with PPO-based reinforcement learning adapted to discrete flows, which enables gradient-guided refinement under shaped reward signals. Importantly, the combination of the two methodologies is simplified by our model’s ability to accept full sequences as input. Note, that for this experiment we used the standard discrete flow sampling from Eq. 2, a short explanation for this choice is given in App. B.3.1.

Genetic Algorithm Because our model operates on full-length sequences, the initial state $x_{t=0}$ is constructed as a mixture of top-performing molecules. We maintain a vocabulary of high-scoring molecules with pairwise Tanimoto distance ≥ 0.7 on Morgan fingerprints. To produce an offspring, two parents are sampled without replacement using rank-based probabilities $p(m) = 1/(\text{rank}(m) + \kappa M)$, where M is the vocabulary size and κ controls the contribution of lower-ranked entries. Parents are then decomposed by the fragmentation rule, and offspring are formed by replacing one fragment of the first parent with a fragment from the second, joined by a separator token in the fragmented SMILES space (i.e. simple string concatenation). By contrast, GenMol relies on predefined reaction templates to form the offspring, which may restrict exploration of chemical space. Naturally, most resulting offspring are not valid molecules, but they only provide the starting state $x_{t=0}$ for our model. While in traditional GA the offspring is mutated, we adopt the mutation operators of Jensen (2019) and apply them to the best-performing molecules found so far to explore their neighborhood.

Reinforcement Learning We adapt PPO (Schulman et al., 2017) to fine-tune the discrete flow policy. Unlike autoregressive models, where the policy log-probability is directly available as the sum of next-token log-likelihoods, discrete flows do not yield a tractable $\log p(x)$ over entire sequences. Instead, following the discrete flow matching framework of Gat et al. (2024), we approximate the log-probability via Monte Carlo estimation over perturbed states. For every sequence, we draw timesteps $t \sim U(0, 1)$, apply the noise schedule to obtain a partially noised state x_t and optimize the time-weighted loss

$$L = \frac{1}{1 - t^2} \sum_{\text{noised positions}} \log \pi_{\theta}(x_1^i | x_t, t). \quad (7)$$

Our rewards are computed as $A = \frac{r - \bar{r}}{\sigma_r + \epsilon}$, where r is the oracle scores, $\bar{r}$ denotes the batch mean, σ_r the standard deviation and ϵ provides numerical stability. The clipped PPO surrogate $\rho(\theta) = \exp(\log \pi_{\theta} - \log \pi_{\text{old}})$ is optimized in the standard way.

Adaptive Sequence Length Sampling As mentioned previously, our model requires a sequence length as input during generation. To accelerate convergence to well-performing molecule lengths, we employ an adaptive bandit that favors lengths with consistently high rewards while still retaining exploration through the prior distribution. We use a peak-seeking variant that blends best-so-far performance, reward quantiles, and an exploration bonus (see Appendix A.3.1).

Combined Optimization Algorithm The overall procedure, combining GA exploration with PPO fine-tuning of the discrete flow, is summarized in Alg. 1, and additional implementation details are provided in Appendix A.3. Unlike GenMol or f -RAG, our method uses a single hyperparameter configuration across all tasks, highlighting that the performance gains arise from the algorithmic design rather than extensive hyperparameter tuning.

Results The PMO benchmark comprises 23 single-objective molecular optimization tasks. Evaluation considers the achieved score and sample efficiency, summarized by the **top10 AUC** metric: the area under the curve of the mean score of the top ten molecules as a function of oracle calls. The scores are normalized to $[0, 1]$, and each run is limited to 10,000 oracle evaluations. GenMol (Lee et al., 2025) and f -RAG (Lee et al., 2024) initialize their populations by screening the entire ZINC250k dataset, i.e. performing 250,000 additional oracle calls before optimization begins. Thus, while they nominally report results with 10k oracle calls, the effective budget is closer to 260,000. To ensure comparability, we evaluate our method under both setups: with prescreening on ZINC250k (Tab.2), directly comparable to GenMol and f -RAG, and without prescreening (Tab.3), directly comparable to baselines such as REINVENT (Olivecrona et al., 2017), MolGA (Tripp and Hernández-Lobato, 2023), and Genetic GFN (Kim et al., 2024), which do not use any prior oracle information. Due

Algorithm 1 Target-Property Optimization (GA + PPO)

Require: Model π , frozen prior π_{old} , oracle O , population $\mathcal{P}$, bandit B , maximum oracle calls N_{max} , fragmentation rule $\mathcal{R}$, mutation op, PPO params $(\epsilon, c_{\text{neg}}, \beta)$

while oracle calls $< N_{\text{max}}$ **do**
 sample parent pairs from $\mathcal{P}$; draw lengths $\ell \sim B$
 $\mathcal{P}_{\text{off}} \leftarrow \{X \sim \pi(\cdot \mid \text{crossover}(\mathcal{R}(p_1), \mathcal{R}(p_2)), \ell)\}$
 $\mathcal{P}_{\text{off}} \leftarrow \mathcal{P}_{\text{off}} \cup \{\text{mutate}(x) : x \in \text{top}_N(\mathcal{P})\}$
 $\mathbf{r} \leftarrow O(\mathcal{P}_{\text{off}})$
 for $k = 1$ to K **do**
 $\theta \leftarrow \theta - \nabla_{\theta} L(\theta; \mathbf{r}, \pi_{\theta}, \pi_{\text{old}})$
 end for
 update B with $(\ell, \mathbf{r})$; $\mathcal{P} \leftarrow \text{top}\{\mathcal{P} \cup \mathcal{P}_{\text{off}}\}$
 $\pi_{\text{old}} \leftarrow \pi$
end while
return top molecules

Table 2: Comparison of models on the PMO benchmark that screen ZINC250k before initialization. We report the AUC-top10 scores, averaged over three runs with standard deviations. Best results and those within one standard deviation of the best are indicated in bold. The scores for f -RAG (Lee et al., 2024) and GenMol (Lee et al., 2025) are taken from the respective publications.

Oracle	InVirtuoGen	GenMol	f-RAG
albuterol similarity	0.975 (± 0.016)	0.937 (± 0.010)	0.977 (± 0.002)
amlodipine mpo	0.836 (± 0.031)	0.810 (± 0.012)	0.749 (± 0.019)
celecoxib rediscovery	0.839 (± 0.013)	0.826 (± 0.018)	0.778 (± 0.007)
deco hop	0.968 (± 0.012)	0.960 (± 0.010)	0.936 (± 0.011)
drd2	0.995 (± 0.000)	0.995 (± 0.000)	0.992 (± 0.000)
fexofenadine mpo	0.904 (± 0.000)	0.894 (± 0.028)	0.856 (± 0.016)
gsk3b	0.988 (± 0.001)	0.986 (± 0.003)	0.969 (± 0.003)
isomers c7h8n2o2	0.988 (± 0.002)	0.942 (± 0.004)	0.955 (± 0.008)
isomers c9h10n2o2pf2cl	0.898 (± 0.018)	0.833 (± 0.014)	0.850 (± 0.005)
jnk3	0.898 (± 0.031)	0.906 (± 0.023)	0.904 (± 0.004)
median1	0.386 (± 0.003)	0.398 (± 0.000)	0.340 (± 0.007)
median2	0.377 (± 0.006)	0.359 (± 0.004)	0.323 (± 0.005)
mestranol similarity	0.991 (± 0.002)	0.982 (± 0.000)	0.671 (± 0.021)
osimertinib mpo	0.881 (± 0.012)	0.876 (± 0.008)	0.866 (± 0.009)
perindopril mpo	0.753 (± 0.019)	0.718 (± 0.012)	0.681 (± 0.017)
qed	0.943 (± 0.000)	0.942 (± 0.000)	0.939 (± 0.001)
ranolazine mpo	0.854 (± 0.012)	0.821 (± 0.011)	0.820 (± 0.016)
scaffold hop	0.711 (± 0.081)	0.628 (± 0.008)	0.576 (± 0.014)
sitagliptin mpo	0.743 (± 0.022)	0.584 (± 0.034)	0.601 (± 0.011)
thiothixene rediscovery	0.652 (± 0.024)	0.692 (± 0.123)	0.584 (± 0.009)
troglitazone rediscovery	0.853 (± 0.003)	0.867 (± 0.022)	0.448 (± 0.017)
valsartan smarts	0.935 (± 0.012)	0.822 (± 0.042)	0.627 (± 0.058)
zaleplon mpo	0.624 (± 0.040)	0.584 (± 0.011)	0.486 (± 0.004)
Sum	18.993 (± 0.219)	18.362	16.928

to space constraints, we only include up to three top-performing baselines, ranked by the sum of AUC-top10 scores across all benchmark tasks. In both setups, InVirtuoGen consistently achieves the best overall performance, considering the sum over all tasks. Ablations in Appendix B.3.1 show that all components of our optimization stack matter. In particular, PPO without prescreening and any genetic algorithm yields higher performance than REINFORCE, a common baseline used in industry.

Table 3: The results of the best performing models on the PMO benchmark, where we quote the AUC-top10 averaged over 3 runs with standard deviations. The best results are highlighted in bold. Values within one standard deviation of the best are also marked in bold. The results for Genetic GFN (Kim et al., 2024) and Mol GA (Tripp and Hernández-Lobato, 2023) are taken from the respective papers. The other results are taken from the original PMO benchmark paper by (Gao et al., 2022a).

Oracle	InVirtuoGen (no prescreen)	Gen. GFN	Mol GA	REINVENT
albuterol similarity	0.950 (± 0.017)	0.949 (± 0.010)	0.896 (± 0.035)	0.882 (± 0.006)
amlodipine mpo	0.733 (± 0.043)	0.761 (± 0.019)	0.688 (± 0.039)	0.635 (± 0.035)
celecoxib rediscovery	0.798 (± 0.028)	0.802 (± 0.029)	0.567 (± 0.083)	0.713 (± 0.067)
deco hop	0.748 (± 0.109)	0.733 (± 0.109)	0.649 (± 0.025)	0.666 (± 0.044)
drd2	0.985 (± 0.002)	0.974 (± 0.006)	0.936 (± 0.016)	0.945 (± 0.007)
fexofenadine mpo	0.845 (± 0.016)	0.856 (± 0.039)	0.825 (± 0.019)	0.784 (± 0.006)
gsk3b	0.952 (± 0.016)	0.881 (± 0.042)	0.843 (± 0.039)	0.865 (± 0.043)
isomers c7h8n2o2	0.968 (± 0.005)	0.969 (± 0.003)	0.878 (± 0.026)	0.852 (± 0.036)
isomers c9h10n2o2pf2cl	0.874 (± 0.013)	0.897 (± 0.007)	0.865 (± 0.012)	0.642 (± 0.054)
jnk3	0.825 (± 0.016)	0.764 (± 0.069)	0.702 (± 0.123)	0.783 (± 0.023)
median1	0.342 (± 0.008)	0.379 (± 0.010)	0.257 (± 0.009)	0.356 (± 0.009)
median2	0.288 (± 0.008)	0.294 (± 0.007)	0.301 (± 0.021)	0.276 (± 0.008)
mestranol similarity	0.797 (± 0.033)	0.708 (± 0.057)	0.591 (± 0.053)	0.618 (± 0.048)
osimertinib mpo	0.870 (± 0.005)	0.860 (± 0.008)	0.844 (± 0.015)	0.837 (± 0.009)
perindopril mpo	0.645 (± 0.032)	0.595 (± 0.014)	0.547 (± 0.022)	0.537 (± 0.016)
qed	0.942 (± 0.000)	0.942 (± 0.000)	0.941 (± 0.001)	0.941 (± 0.000)
ranolazine mpo	0.848 (± 0.010)	0.819 (± 0.018)	0.804 (± 0.011)	0.760 (± 0.009)
scaffold hop	0.589 (± 0.021)	0.615 (± 0.100)	0.527 (± 0.025)	0.560 (± 0.019)
sitagliptin mpo	0.709 (± 0.029)	0.634 (± 0.039)	0.582 (± 0.040)	0.021 (± 0.003)
thiothixene rediscovery	0.625 (± 0.014)	0.583 (± 0.034)	0.519 (± 0.041)	0.534 (± 0.013)
troglitazone rediscovery	0.595 (± 0.053)	0.511 (± 0.054)	0.427 (± 0.031)	0.441 (± 0.032)
valsartan smarts	0.210 (± 0.297)	0.135 (± 0.271)	0.000 (± 0.000)	0.178 (± 0.358)
zaleplon mpo	0.536 (± 0.006)	0.552 (± 0.033)	0.519 (± 0.029)	0.358 (± 0.062)
Sum	16.676 (± 0.256)	16.213	14.708	14.184

3.4 LEAD OPTIMIZATION

Given an initial seed molecule, the goal in lead optimization is to generate leads that exhibit improved binding affinity to a target protein while satisfying constraints on the generated molecules. For the following experiments, the constraints are $\text{QED} \geq 0.6$, $\text{SA} \leq 4$, and Tanimoto similarity $\geq \delta$ to the seed molecule, where $\delta \in \{0.4, 0.6\}$. We follow the benchmark of Lee et al. (2025) and evaluate on five target proteins (parp1, fa7, 5ht1b, braf, jak2), each with three active ligands as molecule seeds. Performance is measured by the docking score (DS) of the most optimized lead (lower is better). We use the same optimization method as in the previous section, but we scale the docking score DS by constraint satisfaction:

$$S(m) = \frac{DS}{15} (1 - \text{penalty}(\text{QED}(m), \text{SA}(m), \text{SIM}(m))). \quad (8)$$

where the penalty increases when $\text{QED} < 0.6$, $\text{SA} > 4$, or $\text{SIM} < \delta$.

Results Table 4 compares our approach against GenMol, RetMol, and GraphGA. Our model consistently achieves competitive or superior docking scores across most proteins and similarity thresholds. Notably, our method remains effective even under the stricter $\delta = 0.6$ constraint, where baseline methods frequently fail to produce improved leads. For example, on parp1 and jak2, our model obtains substantially better docking scores than the baselines. The use of a soft-constraint oracle during training proves advantageous, allowing our model to explore chemical space more effectively while still converging to molecules that meet all constraints. While Tanimoto similarity is a standard proxy for structural similarity, it has known limitations as it reduces complex molecular relationships to a single fingerprint overlap score. To give a more complete view, we also report results without this constraint in Appendix B.4.

Table 4: Docking scores (lower is better) averaged over 3 random seeds. Bold indicates the best result per seed molecule. Values in parentheses indicate solutions with QED > 0.6 and SA < 4 that do not improve the docking score over the seed. For each seed molecule, its docking score, the quantitative estimate of drug-likeness and synthetic accessibility is given.

Protein (DS/QED/SA)	$\delta = 0.4$				$\delta = 0.6$			
	GenMol	RetMol	GraphGA	InVirtuoGen	GenMol	RetMol	GraphGA	InVirtuoGen
parp1								
-7.3/0.888/2.61	-10.6	-9.0	-8.3	-14.1 (± 0.4)	-10.4	-	-8.6	-12.3 (± 0.2)
-7.8/0.758/2.74	-11.0	-10.7	-8.9	-13.4 (± 0.6)	-9.7	-	-8.1	-11.7 (± 0.5)
-8.2/0.438/2.91	-11.3	-10.9	-	-9.0 (± 1.3)	-9.2	-	-	-10.7 (± 0.9)
fa7								
-6.4/0.284/2.29	-8.4	-8.0	-7.8	-8.4 (± 0.4)	-7.3	-7.6	-7.6	-7.7 (± 0.4)
-6.7/0.186/3.39	-8.4	-	-8.2	-8.9 (± 0.5)	-7.6	-	-7.6	-7.5 (± 0.3)
-8.5/0.156/2.66	-	-	-	(-8.0 (± 0.3))	-	-	-	(-7.4 (± 0.4))
5ht1b								
-4.5/0.438/3.93	-12.9	-12.1	-11.7	-13.3 (± 0.1)	-12.1	-	-11.3	-12.4 (± 0.5)
-7.6/0.767/3.29	-12.3	-9.0	-12.1	-12.0 (± 0.7)	-12.0	-10.0	-12.0	-12.0 (± 0.4)
-9.8/0.716/4.69	-11.6	-	-	-10.9 (± 0.2)	-10.5	-	-	-10.6 (± 0.1)
braf								
-9.3/0.235/2.69	-10.8	-11.6	-9.8	-10.1 (± 0.0)	-	-	-	-9.7 (± 0.1)
-9.4/0.346/2.49	-10.8	-	-	-10.8 (± 0.1)	-9.7	-	-	-10.4 (± 0.3)
-9.8/0.255/2.38	-10.6	-	-11.6	-10.6 (± 0.4)	-10.5	-	-10.4	-10.3 (± 0.3)
jak2								
-7.7/0.725/2.89	-10.2	-8.2	-8.7	-10.2 (± 0.8)	-9.3	-8.1	-	-9.7 (± 0.3)
-8.0/0.712/3.09	-10.0	-9.0	-9.2	-10.5 (± 0.3)	-9.4	-	-9.2	-10.4 (± 0.1)
-8.6/0.482/3.10	-9.8	-	-	-10.2 (± 0.2)	-	-	-	-10.3 (± 0.2)
Sum	-148.7	-88.5	-96.3	-152.4 (-160.4)	-117.7	-25.7	-74.8	-145.7 (-153.1)

4 CONCLUSION

We have presented InVirtuoGen, a discrete flow-based generative model for fragmented SMILES. It is a versatile model employable in various stages of common practical drug discovery tasks allowing fragment-level control. By decoupling sequence length from token generation, we can show that a finer granularity during the simulation trajectory leads to more diverse and drug-like molecules in *de novo* generation. The uniform-source formulation also enables seamless integration with a genetic algorithm and PPO-motivated fine-tuning. Across *de novo* generation, fragment-constrained design, and target property optimization, our approach advances the state-of-the-art, achieving a new pareto frontier in quality-diversity trade-offs, with competitive quality and diversity in fragment-constrained tasks, a higher sum of top10 AUC in the PMO benchmark and similarly better results in lead optimization, where docking scores are optimized under constraints.

5 LIMITATIONS

Our fragmented SMILES representation discards stereochemistry, preventing modeling of stereospecific interactions. The rBRICS decomposition may miss chemically relevant fragmentation patterns, and the SA/QED metrics are coarse heuristics that poorly correlate with actual drug-likeness (Skoraczynski et al., 2023; Chen and Jung, 2024). Missing ADMET assessments limit our conclusions and all results remain proxy-based, requiring experimental validation. However, these issues are shared by all compared baselines.

Our sampling modification (Eq. 6) lacks theoretical justification despite strong empirical evidence. For fragment-constrained generation, we employ naive overwriting that disrupts the learned flow dynamics by forcing certain positions to remain fixed throughout the trajectory, contradicting the refinement paradigm central to our approach.

REFERENCES

- P. Auer, N. Cesa-Bianchi, and P. Fischer. Finite-time analysis of the multiarmed bandit problem. *Machine Learning*, 47:235–256, 05 2002. doi: 10.1023/a:1013689704352.
- R. Bickerton, G. Paolini, J. Besnard, S. Muresan, and A. Hopkins. Quantifying the chemical beauty of drugs. *Nature chemistry*, 4:90–8, 02 2012. doi: 10.1038/nchem.1243.
- A. Bou, M. Thomas, S. Dittert, C. N. Ramírez, M. Majewski, Y. Wang, S. Patel, G. Tresadern, M. Ahmad, V. Moens, W. Sherman, S. Sciabola, and G. D. Fabritiis. Acegen: Reinforcement learning of generative chemical agents for drug discovery, 2024. URL <https://arxiv.org/abs/2405.04657>.
- A. Campbell, J. Yim, R. Barzilay, T. Rainforth, and T. Jaakkola. Generative flows on discrete state-spaces: Enabling multimodal flows with applications to protein co-design, 2024. URL <https://arxiv.org/abs/2402.04997>.
- J. Chambers, M. Davies, A. Gaulton, A. Hersey, S. Velankar, R. Petryszak, J. Hastings, L. Bellis, S. McGlinchey, and J. P. Overington. Unichem: a unified chemical structure cross-referencing and identifier tracking system. *Journal of Cheminformatics*, 5(1):3, 2013. doi: 10.1186/1758-2946-5-3. URL <https://doi.org/10.1186/1758-2946-5-3>.
- D. Chen, M. Krimmel, and K. Borgwardt. Flatten graphs as sequences: Transformers are scalable graph generators, 2025. URL <https://arxiv.org/abs/2502.02216>.
- S. Chen and Y. Jung. Estimating the synthetic accessibility of molecules with building block and reaction-aware sascore. *Journal of Cheminformatics*, 16(1):83, 2024. doi: 10.1186/s13321-024-00879-0. URL <https://doi.org/10.1186/s13321-024-00879-0>.
- Y. David and N. Shimkin. Pure exploration for max-quantile bandits. volume 9851, pages 556–571, 09 2016. ISBN 978-3-319-46127-4. doi: 10.1007/978-3-319-46128-1_35.
- J. Degen, C. Wegscheid-Gerlach, A. Zaliani, and M. Rarey. On the art of compiling and using ‘drug-like’ chemical fragment spaces. *ChemMedChem: Chemistry Enabling Drug Discovery*, 3(10):1503–1507, 2008.
- Y. El Mesbahi and E. Noutahi. Safe setup for generative molecular design. *arXiv preprint arXiv:2410.20232*, 2024. preprint.
- P. Ertl and A. Schuffenhauer. Estimation of synthetic accessibility score of drug-like molecules based on molecular complexity and fragment contributions. *Journal of cheminformatics*, 1:8, 06 2009. doi: 10.1186/1758-2946-1-8.
- W. Gao, T. Fu, J. Sun, and C. W. Coley. Sample efficiency matters: A benchmark for practical molecular optimization, 2022a. URL <https://arxiv.org/abs/2206.12411>.
- W. Gao, R. Mercado, and C. W. Coley. Amortized tree generation for bottom-up synthesis planning and synthesizable molecular design, 2022b. URL <https://arxiv.org/abs/2110.06389>.
- I. Gat, T. Remez, N. Shaul, F. Kreuk, R. T. Q. Chen, G. Synnaeve, Y. Adi, and Y. Lipman. Discrete flow matching. *arXiv preprint arXiv:2407.15595*, 2024.
- J. H. Jensen. A graph-based genetic algorithm and generative model/monte carlo tree search for the exploration of chemical space. *Chem. Sci.*, 10:3567–3572, 2019. doi: 10.1039/c8sc05372c. URL <http://dx.doi.org/10.1039/C8SC05372C>.
- W. Jin, R. Barzilay, and T. Jaakkola. Junction tree variational autoencoder for molecular graph generation. *Proceedings of the 35th International Conference on Machine Learning*, 80:2323–2332, 2018.
- W. Jin, D. Barzilay, and T. Jaakkola. Multi-objective molecule generation using interpretable substructures. In H. D. III and A. Singh, editors, *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pages 4849–4859. Pmlr, 13–18 Jul 2020. URL <https://proceedings.mlr.press/v119/jin20b.html>.

- S. Jinsong, J. Qifeng, C. Xing, Y. Hao, and L. Wang. Molecular fragmentation as a crucial step in the ai-based drug development pathway. *Communications Chemistry*, 7(1):20, 2024. doi: 10.1038/s42004-024-01109-2. URL <https://doi.org/10.1038/s42004-024-01109-2>.
- H. Kim, M. Kim, S. Choi, and J. Park. Genetic-guided gflownets for sample efficient molecular optimization, 2024. URL <https://arxiv.org/abs/2402.05961>.
- T. Kirkman, C. Silva, M. Tosin, and M. Dias. How to find a fragment: Methods for screening and validation in fragment-based drug discovery. *ChemMedChem*, 19, Nov. 2024. doi: 10.1002/cmdc.202400342.
- P. Kirsch, A. M. Hartman, A. K. H. Hirsch, and M. Empting. Concepts and core principles of fragment-based drug design. *Molecules*, 24(23), 2019. ISSN 1420-3049. doi: 10.3390/molecules24234309. URL <https://www.mdpi.com/1420-3049/24/23/4309>.
- S. Lee, K. Kreis, S. P. Veccham, M. Liu, D. Reidenbach, S. Paliwal, A. Vahdat, and W. Nie. Molecule generation with fragment retrieval augmentation, 2024. URL <https://arxiv.org/abs/2411.12078>.
- S. Lee, K. Kreis, S. P. Veccham, M. Liu, D. Reidenbach, Y. Peng, S. Paliwal, W. Nie, and A. Vahdat. Genmol: A drug discovery generalist with discrete diffusion, 2025. URL <https://arxiv.org/abs/2501.06158>.
- I. Loshchilov and F. Hutter. Decoupled weight decay regularization, 2019. URL <https://arxiv.org/abs/1711.05101>.
- K. Maziarz, H. Jackson-Flux, P. Cameron, F. Sirockin, N. Schneider, N. Stiefl, M. Segler, and M. Brockschmidt. Learning to extend molecular scaffolds with structural motifs, 2024. URL <https://arxiv.org/abs/2103.03864>.
- O. M. Morrison, F. Pichi, and J. S. Hesthaven. Gfn: A graph feedforward network for resolution-invariant reduced operator learning in multifidelity applications, 2024. URL <https://arxiv.org/abs/2406.03569>.
- A. Nigam, P. Friederich, M. Krenn, and A. Aspuru-Guzik. Augmenting genetic algorithms with deep neural networks for exploring the chemical space, 2020. URL <https://arxiv.org/abs/1909.11655>.
- E. Noutahi, C. Gabellini, M. Craig, J. S. C. Lim, and P. Tossou. Gotta be safe: A new framework for molecular design, 2023. URL <https://arxiv.org/abs/2310.10773>.
- M. Olivecrona, T. Blaschke, O. Engkvist, and H. Chen. Molecular de-novo design through deep reinforcement learning. *Journal of Cheminformatics*, 9(1):1–14, 2017.
- W. Peebles and S. Xie. Scalable diffusion models with transformers. *arXiv preprint arXiv:2212.09748*, 2022.
- D. Polykovskiy, A. Zhebrak, B. Sanchez-Lengeling, S. Golovanov, O. Tatanov, S. Belyaev, R. Kurbanov, A. Artamonov, V. Aladinskiy, M. Veselov, A. Kadurin, S. Johansson, H. Chen, S. Nikolenko, A. Aspuru-Guzik, and A. Zhavoronkov. Molecular sets (moses): A benchmarking platform for molecular generation models, 2018.
- D. Rogers and M. Hahn. Extended-connectivity fingerprints. *Journal of Chemical Information and Modeling*, 50(5):742–754, 05 2010. doi: 10.1021/ci100050t. URL <https://doi.org/10.1021/ci100050t>.
- S. S. Sahoo, M. Arriola, Y. Schiff, A. Gokaslan, E. Marroquin, J. T. Chiu, A. Rush, and V. Kuleshov. Simple and effective masked diffusion language models, 2024. URL <https://arxiv.org/abs/2406.07524>.
- Y. Schiff, S. S. Sahoo, H. Phung, G. Wang, S. Boshar, H. Dalla-torre, B. P. de Almeida, A. Rush, T. Pierrot, and V. Kuleshov. Simple guidance mechanisms for discrete diffusion models, 2025. URL <https://arxiv.org/abs/2412.10193>.

- J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- G. Skoraczynski, M. Kitlas, B. Miasojedow, and A. Gambin. Critical assessment of synthetic accessibility scores in computer-assisted synthesis planning. *Journal of Cheminformatics*, 15(1):6, 2023. doi: 10.1186/s13321-023-00678-z. URL <https://doi.org/10.1186/s13321-023-00678-z>.
- T. Sterling and J. Irwin. Zinc 15 - ligand discovery for everyone. *Journal of chemical information and modeling*, 55, Oct. 2015. doi: 10.1021/acs.jcim.5b00559.
- J. Su, Y. Lu, S. Pan, A. Murtadha, B. Wen, and Y. Liu. Roformer: Enhanced transformer with rotary position embedding, 2023. URL <https://arxiv.org/abs/2104.09864>.
- A. Tripp and J. M. Hernández-Lobato. Genetic algorithms are strong baselines for molecule generation, 2023. URL <https://arxiv.org/abs/2310.09267>.
- D. Weininger. Smiles, a chemical language and information system. 1. introduction to methodology and encoding rules. *Journal of Chemical Information and Computer Sciences*, 28(1):31–36, 02 1988. doi: 10.1021/ci00057a005. URL <https://doi.org/10.1021/ci00057a005>.
- Y. Xie, C. Shi, H. Zhou, Y. Yang, W. Zhang, Y. Yu, and L. Li. Mars: Markov molecular sampling for multi-objective drug discovery, 2021. URL <https://arxiv.org/abs/2103.10432>.
- X. Zeng, F. Wang, Y. Luo, S. gu Kang, J. Tang, F. C. Lightstone, E. F. Fang, W. Cornell, R. Nussinov, and F. Cheng. Deep generative molecular design reshapes drug discovery. *Cell Reports Medicine*, 3(12):100794, 2022.

A EXPERIMENTAL DETAILS

A.1 IMPLEMENTATION DETAILS

The backbone model, adapted from Peebles and Xie (2022), has 36 layers⁶, 12 heads, the hidden dimension is chosen as 768, and uses rotary positional embeddings (Su et al., 2023). Training is performed for one epoch with batch size of 300. Sequences of similar length are bucketed together to reduce the padding per batch. The maximal number of tokens per batch is limited to 25,000. The AdamW (Loshchilov and Hutter, 2019) optimizer, with $(\beta_1 = 0.99, \beta_2 = 0.999)$ is used with a learning rate of 10^{-4} . Additionally, the learning rate is varied according to a linear warm-up cosine annealing scheduler.

A.2 FRAGMENT-CONSTRAINED GENERATION PARAMETERS

The time-granularity of 0.01 is used together with our proposed sampling from Eq. 6. We see a slight performance increase when increasing the sampling time-granularity, as demonstrated in Appendix B.2.

A.3 TARGET-PROPERTY OPTIMIZATION PARAMETERS

A smaller model is used in this setting, with 12 layers instead of 36. The population for the genetic algorithm, which provides the starting points of the trajectory $\{x_{t=0}\}$, is updated every 50 scored samples. Out of 80 generated samples, 20 are obtained by mutating the top 20 candidates generated so far, using the mutation operators from Jensen (2019). For rank-based sampling, we set $\kappa = 0.001$. For the calculation of the Tanimoto distance, we use Morgan fingerprints with 2048 bits and a radius of 2 nodes. The model is updated with PPO after every 100 scored SMILES. For each sequence, 50 timesteps $t \sim U(0, 1)$ are sampled to construct the training dataset $\{x_t\}$ by interpolating between a sequence of random tokens and the sampled values. During each epoch, 10 optimization steps are performed with a learning rate of 10^{-4} . The clipping coefficient for PPO is set to $\eta = 0.2$.

When constructing the initial population by prescreening ZINC250K, no experience replay is used. Otherwise, an experience replay buffer of 300 samples is maintained.

A.3.1 PEAK-SEEKER BANDIT

Our peak-seeking bandit, given in Alg. 2 is a heuristic that combines elements of UCB bandits (Auer et al., 2002) with quantile-based bandit updates (David and Shimkin, 2016).

Algorithm 2 Peak-Seeker Bandit for Adaptive Length Sampling

Require: Candidate lengths $\{n_k\}$, prior $\pi^{(0)}$, quantile target q , learning rate η_q , weights $(w_{\text{best}}, w_{\text{quant}})$, bandwidth σ , exploration bonus c , temperature τ , floor probability ϵ

- 1: **while** generating molecules **do**
- 2: For each arm k , track visit count N_k , best reward b_k , and running quantile estimate $\hat{q}_k$
- 3: Compute score $s_k \leftarrow w_{\text{best}}b_k + w_{\text{quant}}\hat{q}_k + \text{UCB}(N_k) + \text{neighborhood}(L_{\text{best}}, n_k)$
- 4: For unvisited arms: use prior $\pi^{(0)}$ as fallback
- 5: Convert scores to probabilities $p_k \propto \exp(s_k/\tau)$, apply floor ϵ , normalize
- 6: Sample sequence length $L \sim \text{Categorical}(p)$
- 7: Generate molecules of length L and obtain reward r
- 8: Update $N_k, b_k \leftarrow \max(b_k, r), \hat{q}_k$ by quantile SGD
- 9: Update global best $(L_{\text{best}}, r_{\text{best}})$ if $r > r_{\text{best}}$
- 10: **end while**

⁶For the results in the fragment-constrained generation and target property/lead optimization a smaller model with 12 layers was used to improve speed and reduce memory use.

B ADDITIONAL EXPERIMENTAL RESULTS

B.1 DE NOVO GENERATION

Sequence Length Distribution In Fig. 4 we show the sequence length distribution of ZINC250K. The maximum observed length is 84, which implies that a masked model would require up to 84 steps to produce a sample. In contrast, our discrete flow model with a uniform source decouples the number of steps from the sequence length, allowing shorter sequences to be refined more within the same compute budget.

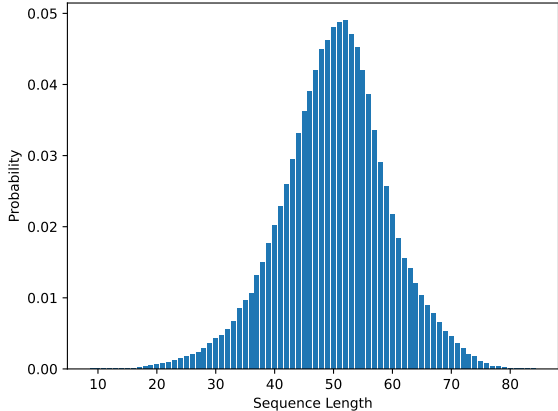


Figure 4: Sequence length distribution of ZINC250K. The maximum observed length is 84, which implies that a masked model requires at least 84 sampling steps, putting it close to our step size $h = 0.01$.

Non-Curated Samples In Fig. 5 we provide non-curated samples from *de novo* generation with temperature $T = 1$ and randomness $r = 0$.

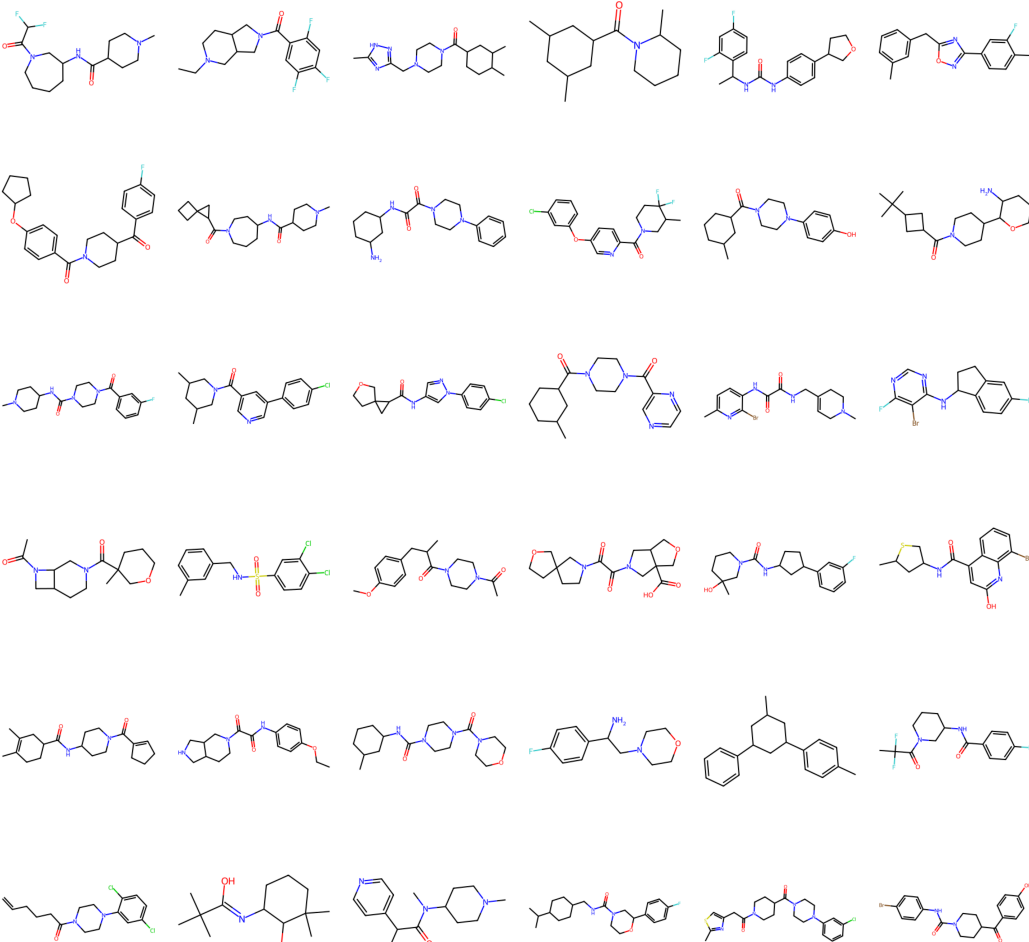


Figure 5: Non curated samples for *de novo* generation

Comparison of Sampling Methods In Fig. 6, we compare the quality versus diversity trade-off generated by sampling according to our proposed sampling method (Eq. 6) and for the one proposed by Gat et al. (2024) (Eq. 2). The benefit of increasing the sampling granularity disappears when standard sampling is used. In the following, we investigate possible causes for this.

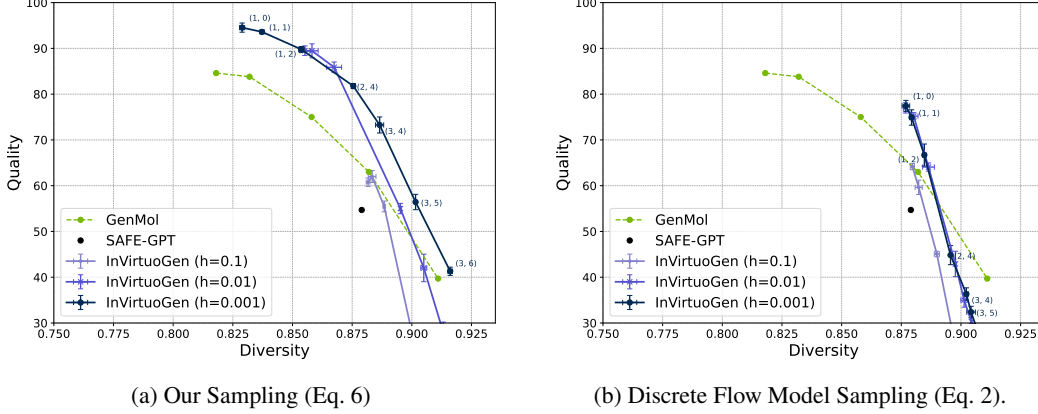
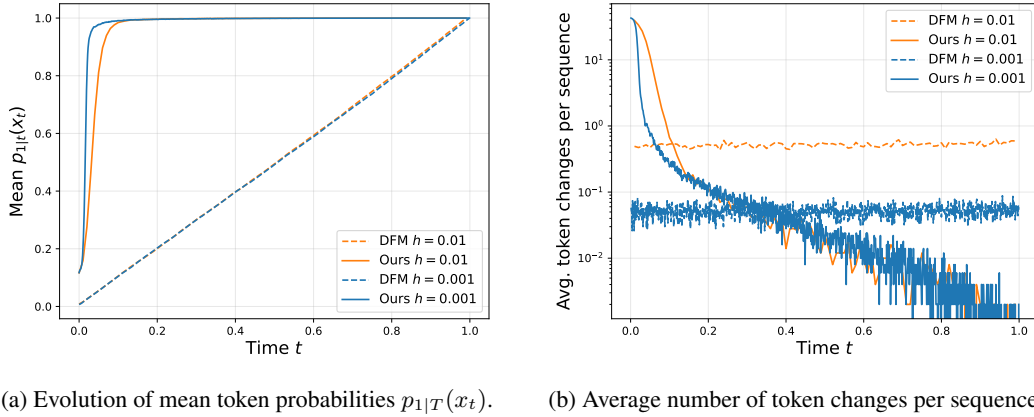


Figure 6: Comparison of the different sampling methods in terms of quality vs diversity.



(a) Evolution of mean token probabilities $p_{1:T}(x_t)$. (b) Average number of token changes per sequence.

Figure 7: Dynamic behavior during sampling. Our method rapidly concentrates token probabilities and progressively reduces the number of token changes, while standard sampling (DFM) increases confidence linearly and keeps changes nearly constant.

We provide statistics aggregated from the sampling trajectories. As shown in Fig. 7a, for our sampling method the probability density of the current tokens $p_{1:T}(x_t)$ quickly saturates at 1, while for standard sampling (Eq.2) it increases almost linearly across steps. This difference is mirrored in the dynamics of token changes (Fig. 7b): our approach begins with many parallel updates that gradually decay into a refinement phase, whereas standard sampling maintains a nearly constant but low rate of changes throughout. Our decay pattern is closer to the intuitive notion of refinement, where the number of modifications decreases as the sequence converges towards a high-probability solution. Finally, the cumulative number of token changes (Fig. 8) highlights a fundamental distinction: DFM sampling yields a step-size-invariant total number of modifications, effectively fixing the update budget, while our method scales with the granularity h , allowing more structured refinements when smaller steps are used. The observation that the number of token changes is independent of time granularity for standard DFM sampling provides a plausible explanation for the lack of improvement in standard sampling at higher resolutions.

Additionally, Figs. 9a and 10b show QED and SA distributions across time resolutions h . Under our sampling (Eq. 6), the QED distribution shifts markedly upward to higher values, while SA also seems to move to smaller values (Figs. 9a, 9b). In contrast, the standard discrete flow update (Eq. 2)

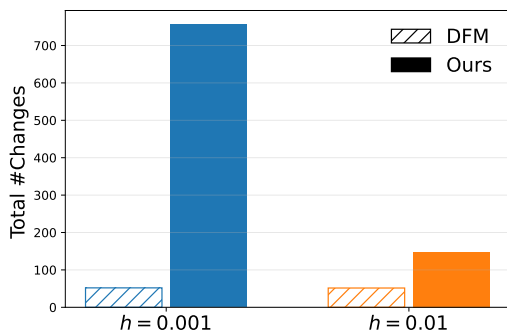


Figure 8: Cumulative number of token changes per sequence for different time resolutions h . DFM sampling shows step-size invariance, while our method scales the number of refinements with stepsize h .

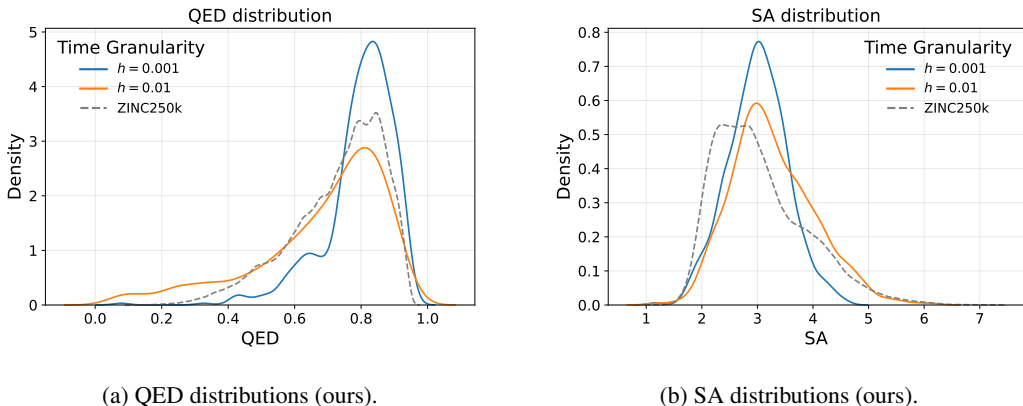


Figure 9: Distributions of QED and SA for molecules generated with our sampling rule (Eq. 6) at different time resolutions h . ZINC250k is shown as a dashed baseline.

more closely follows ZINC250k for both metrics (Figs. 10a, 10b). This upward shift in QED and downward shift in SA under our sampling method is consistent with the stronger quality–diversity frontier reported earlier. Although one might expect the generated distributions to overlap with ZINC250k, it is important to emphasize that ZINC250k contributed only a small fraction of the training data (2.5×10^5 molecules vs. 10^9 in total). The training dataset contains biological compounds, which tend to be longer and complex and therefore exhibit higher synthetic accessibility. Moreover, the objective of our framework is not to reproduce the underlying training distribution but to generate drug-like molecules with optimized properties. Figure 11 compares standard sampling (Eq. 2) and our sampling (Eq. 6) at ($T = 1, r = 0$). The left panel shows the average number of carbon atoms per sequence. Our sampling results in a slightly higher carbon frequency, though the difference is modest.

The right panel reports the distribution of the number of fragments per sequence, estimated from the maximal attachment index observed. Standard sampling produces a higher fraction of sequences with many fragments, whereas our sampling shifts weight toward intermediate fragment counts. We hypothesize that this structural difference translates into improved drug-likeness, but further investigation is necessary.

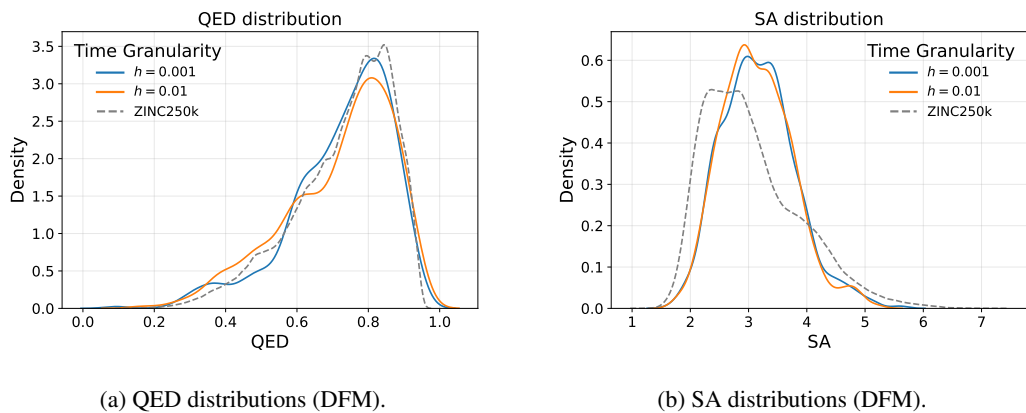


Figure 10: Distributions of QED and SA for molecules generated with the standard discrete flow update (Eq. 2) across different time resolutions h . ZINC250k is shown as a dashed baseline.

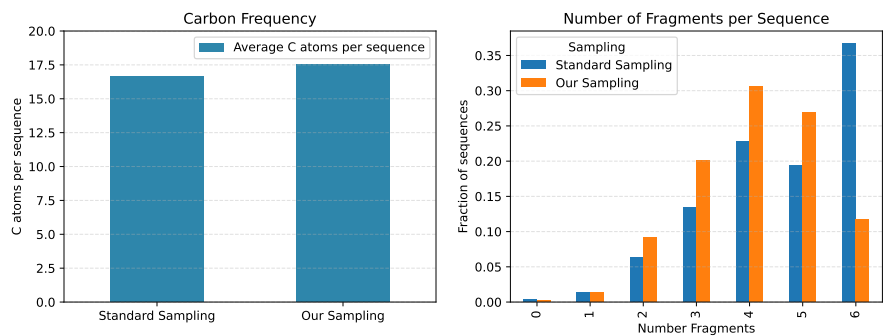


Figure 11: Sampling comparison at temperature $T = 1$ and noise $r = 0$. Left: average number of carbon atoms per sequence. Right: distribution of the number of fragments per sequence.

Impact of Time-Weighting Loss In Fig. 12, we show results obtained by training the same model as for the other *de novo* generation experiments, but without time-weighting the loss terms. While the performance under sampling with Eq. 2 is nearly unchanged, the results with Eq. 6 are significantly worse without time-weighting, showing that our modification has a significant impact on the result.

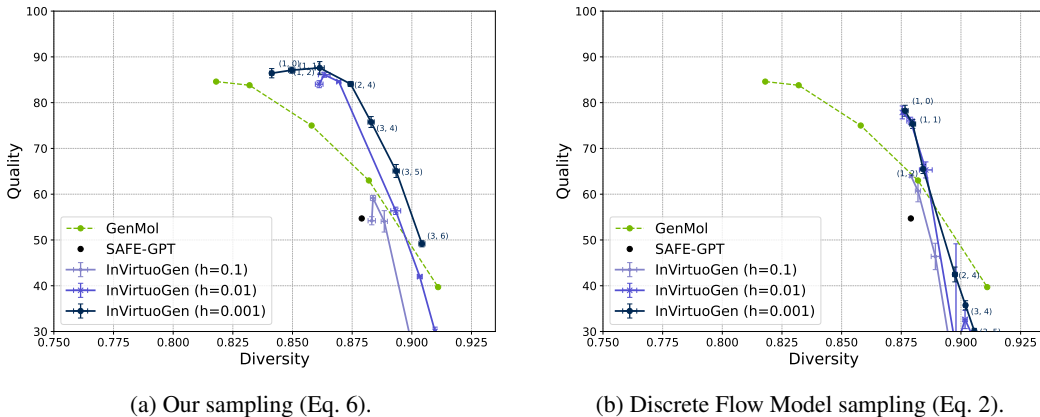


Figure 12: Quality-diversity trade-off without using the time-weighting $\frac{1}{1-t^2}$ of the loss, given in Eq. 4. Left: results from sampling with Eq. 6. Right: results from sampling with Eq. 2.

Timing Studies To compare with GenMol, we also report results for the smaller 12-layer model used in the target-property optimization section. Its quality-diversity frontier is shown in Fig. 13. Sampling with time granularity $h = 0.01$ on an RTX 4090 yields 20.2 ± 0.3 s for 1000 samples. Using the same setup with GenMol (instantiated from the provided configuration rather than a released checkpoint) we obtain 33.2 ± 0.3 s. We emphasize, however, that speed is a minor factor in drug discovery, since downstream steps such as docking or wet-lab experiments typically dominate runtime.

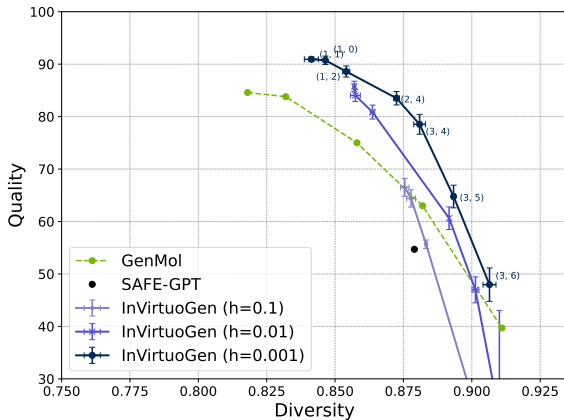
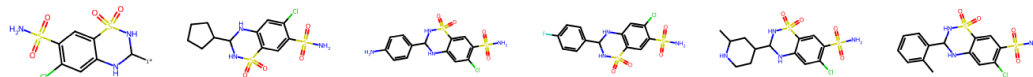


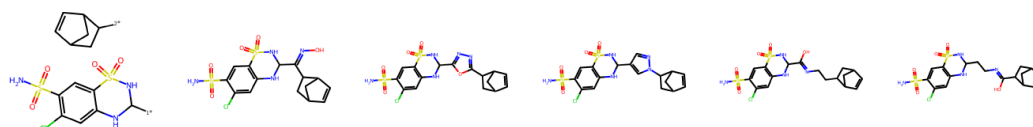
Figure 13: Quality-diversity frontier for the 12-layer model. At $h = 0.01$, InVirtuoGen achieves higher maximum quality to GenMol while attaining substantially higher diversity in the high-quality regime.

B.2 FRAGMENT-CONSTRAINED GENERATION

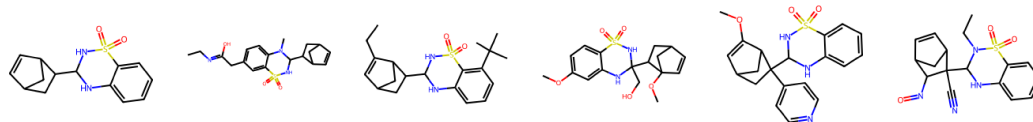
In Fig. 14 we provide non-cherry picked samples for the fragment-constrained generation task. The left-most figure in every row depicts the starting fragment(s).



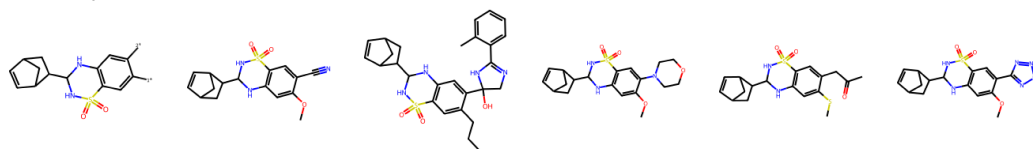
(a) Motif Extension



(b) Linker Design/Scaffold Morphing



(c) Superstructure Generation



(d) Scaffold Decoration

Figure 14: Non cherry-picked samples generated by InVirtuoGen on fragment-constrained design tasks. The left-most figure in every row depicts the starting fragment(s).

B.3 TARGET-PROPERTY OPTIMIZATION

B.3.1 ABLATION STUDIES

In Table 5 we show the results of ablation studies of the core components of our optimization framework. We include results for the following ablations:

- Including the experience replay in the prescreened setting.
- Sampling according to eq. 6 but with start time $t_{start} = 0.2$.
- Sampling according to eq. 6.
- Sampling sequence lengths from the Zinc250K distribution instead of using our Peak-Finder Bandit.
- No mutation applied to the best performing molecules.
- No PPO, relying only on the genetic algorithm.
- No prescreening, but leaving out the experience replay, yielding a significantly lower performance.
- A baseline without GA, mutation, or prescreening, which allows a fair comparison to REINVENT (Olivecrona et al., 2017) and slightly better performance.

As the table shows, the results obtained with our sampling method depend strongly on the start time of the trajectory simulation. We attribute this sensitivity to the dynamics in Fig. 7b: at early times many positions are simultaneously updated, which can diminish the performance gains provided by the genetic algorithm component. And because we did not want to tune our parameters, we stuck to sampling according to Eq 2.

B.4 LEAD OPTIMIZATION

We provide the results obtained when not requiring similarity to the seed molecule in Tab. 6, showing that performance improves significantly when this restriction is removed. Constraints on drug-likeness and synthesizability are still enforced.

Table 5: Ablation study results on the PMO benchmark. We report the AUC-top10 scores from single runs. Best results are highlighted in bold.

Oracle	With Experience Replay	Sampling with eq. 6, $t_{start} = 0.2$	Sampling with eq. 6	No Bandit	No Mutation	No PPO	No Prescreen, No Experience Replay	No Prompter, No Mutation, No Prescreen
albuterol similarity	0.993	0.975	0.972	0.991	0.928	0.881	0.969	0.851
amlodipine mpo	0.813	0.795	0.801	0.764	0.777	0.755	0.679	0.547
celecoxib rediscovery	0.872	0.866	0.795	0.785	0.758	0.815	0.834	0.768
deco hop	0.960	0.965	0.973	0.988	0.970	0.943	0.651	0.659
drd2	0.995	0.995	0.995	0.995	0.986	0.986	0.984	0.950
fexofenadine mpo	0.908	0.870	0.879	0.913	0.834	0.853	0.852	0.735
gsk3b	0.990	0.989	0.987	0.992	0.976	0.972	0.949	0.867
isomers c7h8n2o2	0.989	0.986	0.986	0.978	0.967	0.974	0.971	0.857
isomers c9h10n2o2pf2cl	0.876	0.882	0.871	0.919	0.851	0.834	0.888	0.798
jnk3	0.892	0.917	0.881	0.917	0.807	0.889	0.825	0.701
median1	0.382	0.379	0.370	0.385	0.370	0.363	0.342	0.308
median2	0.333	0.386	0.367	0.377	0.331	0.363	0.304	0.241
mestranol similarity	0.992	0.986	0.986	0.983	0.963	0.962	0.723	0.764
osimertinib mpo	0.878	0.868	0.870	0.868	0.857	0.858	0.864	0.802
perindopril mpo	0.733	0.753	0.696	0.710	0.695	0.682	0.627	0.463
qed	0.943	0.943	0.943	0.944	0.943	0.943	0.943	0.941
ranolazine mpo	0.878	0.840	0.843	0.866	0.807	0.816	0.840	0.762
scaffold hop	0.654	0.632	0.648	0.657	0.790	0.595	0.619	0.522
sitagliptin mpo	0.772	0.766	0.613	0.738	0.457	0.546	0.693	0.324
thiothixene rediscovery	0.685	0.625	0.542	0.625	0.588	0.526	0.620	0.465
troglitazone rediscovery	0.870	0.859	0.812	0.855	0.833	0.836	0.434	0.418
valsartan smarts	0.870	0.906	0.891	0.927	0.766	0.715	0.000	0.150
zaleplon mpo	0.617	0.654	0.643	0.605	0.576	0.653	0.520	0.458
Sum	18.893	18.836	18.364	18.782	17.831	17.758	16.131	14.349

Table 6: Docking scores (lower is better) averaged over 3 seeds. Bold indicates the best result per seed. For each seed molecule, its docking score, the quantitative estimate of drug-likeness and synthetic accessibility is given.

Protein (DS/QED/SA)	No sim
	InVirtuoGen
5ht1b	
-4.5/0.438/3.93	-12.5 (± 1.0)
-7.6/0.767/3.29	-13.4 (± 0.5)
-9.8/0.716/4.69	-13.0 (± 0.3)
braf	
-9.3/0.235/2.69	-12.5 (± 0.2)
-9.4/0.346/2.49	-12.3 (± 0.8)
-9.8/0.255/2.38	-12.2 (± 0.4)
fa7	
-6.4/0.284/2.29	-9.9 (± 0.4)
-6.7/0.186/3.39	-10.2 (± 0.5)
-8.5/0.156/2.66	-9.4 (± 0.2)
jak2	
-7.7/0.725/2.89	-12.0 (± 0.3)
-8.0/0.712/3.09	-12.1 (± 0.2)
-8.6/0.482/3.10	-11.6 (± 0.7)
parp1	
-7.3/0.888/2.61	-13.5 (± 0.2)
-7.8/0.758/2.74	-13.3 (± 0.7)
-8.2/0.438/2.91	-13.5 (± 0.5)
Sum	-181.4