

ACTIVE LEARNING ON SYNTHONS FOR MOLECULAR DESIGN

Tom George Grigg^{1†} Mason Burlage^{1†} Oliver Scott¹ Adam Taouil¹ Dominique Sydow¹
Liam Wilbraham¹

¹Recursion

liam.wilbraham@recursion.com

ABSTRACT

Exhaustive virtual screening is highly informative but often intractable against the expensive objective functions involved in modern drug discovery. This problem is exacerbated in combinatorial contexts such as multi-vector expansion, where molecular spaces can quickly become ultra-large. Here, we introduce Scalable Active Learning via Synthon Acquisition (SALSA): a simple algorithm applicable to multi-vector expansion which extends pool-based active learning to non-enumerable spaces by factoring modeling and acquisition over synthon or fragment choices. Through experiments on ligand- and structure-based objectives, we highlight SALSA’s sample efficiency, and its ability to scale to spaces of trillions of compounds. Further, we demonstrate application toward multi-parameter objective design tasks on three protein targets – finding SALSA-generated molecules have comparable chemical property profiles to known bioactives, and exhibit greater diversity and higher scores over an industry-leading generative approach.

1 INTRODUCTION

Given the strong association between a molecule’s core scaffold and its chemical properties, a common workflow is to iteratively design, make, and test changes at targeted R-groups in order to advance therapeutics through the discovery pipeline (Schneider, 2017). Exhaustive virtual screening of R-group changes aids designers and medicinal chemists in the search for promising, synthesizable molecular structures, but quickly becomes intractable against computationally expensive scores as the number of possible attachments increases. Prior work in MolPAL (Graff et al., 2021) extends the scope of screening to spaces on the order of 100M molecules via pool-based deep Bayesian optimization. However, in the context of multi-vector expansion, where multiple R-groups are explored simultaneously, spaces can easily surpass 100B+ possible combinations in early stage discovery. At this scale, designers often turn to specialized cheminformatics tools which can be configured to screen constrained synthesizable spaces for substructure (Schmidt et al., 2021), similarity (Bellmann et al., 2021; Cheng & Beroza, 2023), and docking-based (Sadybekov et al., 2022) design objectives.

For bespoke or multi-parameter objectives (MPOs), designers may employ generative (or *inverse*) design. Modern generative approaches typically optimize pre-trained prior distributions on graphs or SMILES towards a molecular score e.g. via RL (Loeffler et al., 2024), guided diffusion (Weiss et al., 2023), etc. Historically, these methods faced issues with synthetic accessibility (Gao & Coley, 2020; Renz et al., 2019), but recent works mitigate with explicit synthesis constraints (Grisoni et al., 2021; Bradshaw et al., 2019; Fialkova et al., 2021), or analogizing (Shitong Luo, 2024; Gao et al., 2024). However, their usage remains limited in practice due to the unwieldy tension between synthesizability and drug-likeness versus novelty when sampling from a generative molecular model.

Here, we extend the domain of pool-based active learning (AL) to a multi-vector expansion context by introducing Scalable Active Learning via Synthon Acquisition (SALSA). By factoring learning over independent synthon or fragment choices, SALSA facilitates screening in explicitly configurable molecular spaces on the order of trillions of compounds. We demonstrate that SALSA is sample efficient with respect to baselines, and validate its application to multi-vector design tasks on three protein targets. We find that SALSA identifies molecules with comparable chemical property distributions to known bioactive compounds while optimizing pharmacophore and structure-based MPOs – improving on established generative approaches, and offering a pragmatic alternative.

[†]Equal contribution

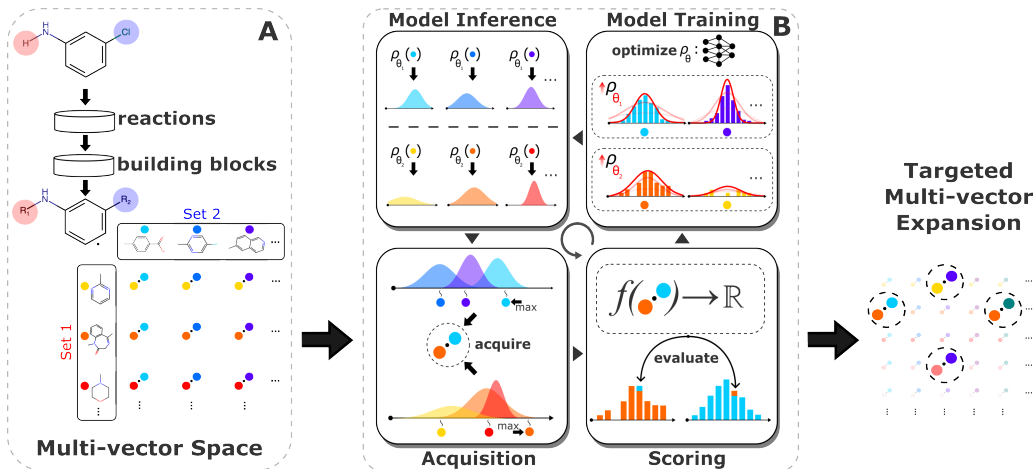


Figure 1: **A** Construction of a 2-vector synthon space. **B** AL loop against scoring function f .

2 METHODS

Search space SALSA consumes as input a target molecular space formed by pre-defined choices of synthons or fragments, as well as a molecular objective function f . Fig. 1A exemplifies construction of a target space for a simple 2-vector expansion scheme on a core with two R-groups. Given a set of SMIRKS-encoded reactions and building blocks, applicable chemistry is represented by a synthon set $\mathcal{S}_i$ for each vector, determined by efficient pattern matching. In our experiments, we use building blocks from Mcule (accessed 18 Sep 2023), and a set of custom SMIRKS (see A.2).

Synthon acquisition Fig. 1B illustrates SALSA for 2-vector expansion. K molecules are initially sampled randomly and scored with f . Scores are recorded for each molecule’s constituent synthons, and a surrogate model is trained at each vector. Score distributions are predicted for all synthons, and K new molecules are sampled via an acquisition strategy α . These molecules are scored to form additional synthon datapoints with which to retrain. This process loops for N rounds, or until convergence – acquiring up to $N \times K$ molecules. Conceptually, SALSA navigates two (or n , in general) non-stationary multi-armed bandit problems simultaneously, one for the pool of synthons at each vector. Factoring the decision problem in this way nullifies inference-time combinatorial complexity from $O(\prod_i |\mathcal{S}_i|)$ to $O(\sum_i |\mathcal{S}_i|)$, which is the primary limitation for full-molecular AL.

In our experiments, we adapt the (approximate) Thompson sampling (TS) strategy outlined in MolPAL (Graff et al., 2021) to the multi-vector case. Here, synthon acquisition scores are sampled from a predicted Gaussian, i.e. $\alpha(s) \sim \mathcal{N}(\mu_\theta(s), \sigma_\theta(s)) \forall s \in \mathcal{S}_i$ for learned parameters θ . Top-scoring synthons are combined and the resulting molecule is scored if unseen, otherwise synthons are resampled. The loop terminates early if more than a given threshold of samples are rejected in a round – as this suggests convergence in acquisition probability. We include pseudo-code describing SALSA in A.1. We also adapted and investigated alternative acquisition strategies (see A.4), as well as a variant that uses one model for all synthon sets (A.5), rather than a model per vector.

Surrogate models We adopt chemprop’s (Heid et al., 2024) implementation of a directed message-passing neural network (MPNN) as our choice of surrogate model. The MPNN dynamically encodes a feature vector by aggregating rounds of message passing across the bonds of a molecule’s 2D graph. A feed-forward head with two output nodes then operates on this graph-based representation to predict a mean $\mu_\theta(s)$ and variance $\sigma_\theta(s)$ for a synthon s . We train these models end-to-end to predict synthon score distributions via a mean-variance estimation (MVE) loss $\mathcal{L}(y, s, \theta) = \frac{\log 2\pi}{2} + \log \sigma_\theta(s) - \frac{1}{2} \left(\frac{y - \mu_\theta(s)}{\sigma_\theta(s)} \right)^2$ – i.e. maximum-likelihood estimation of Gaussian density for observed synthon-score pairs $(s, y) \in \mathcal{S}_i \times \mathbb{R}$ (up to regularization induced priors). We found this model to perform better than fixed-feature alternatives (A.4), consistent with findings in MolPAL. Details on architecture and hyperparameter choices are included in A.3.

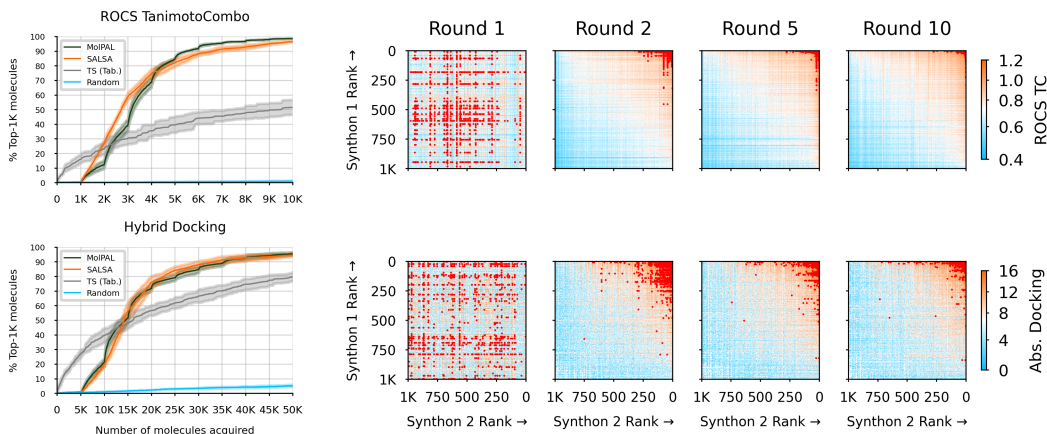


Figure 2: Recall of top-1K compounds in the 1M target space for ROCS-TC (top) and docking (bottom) as a function of molecules acquired, smoothed over 5 trials. The heatmaps show the enumerated target space decomposed across synthon axes and coloured by score. For a given SALSA round, synthons are ordered by Monte Carlo-estimated acquisition probability, i.e. the likelihood of sampling increases moving up and right. Top-1K ground truth molecules are highlighted in red.

3 EXPERIMENTS

Sample efficiency We begin by demonstrating SALSA’s performance in an enumerated 1M molecule space, using CDK2 as a model system (PDB: 6GUH (Wood et al., 2019)). To set-up our expansion, we isolated the co-crystallized ligand’s core and functionalized reaction handles at two vector positions. Synthons were generated at each vector (see A.2), resulting in 910K and 2.4M synthons respectively. 1K synthons were subsampled from each set to create a $1K \times 1K = 1M$ -size space as desired. We defined a docking score using the protein structure, and a 3D shape/color similarity score using the co-crystallized ligand as a reference via OpenEye Hybrid Docking and ROCS TanimotoCombo Score (ROCS-TC), respectively (see A.7). The space was exhaustively enumerated and scored with both objectives to obtain ground truth, enabling comparison of SALSA to: TS (Tab.) (Klarich et al., 2024) which is conceptually similar but updates tabular Gaussian models with exact fixed-variance TS, and a MolPAL-like (Graff et al., 2021) full-molecular pool-based screen.

Fig. 2 depicts 10 rounds of SALSA for each task, with an objective scoring budget of 1K and 5K molecules per round for ROCS-TC and docking, respectively. SALSA is able to ultimately identify 96.5% and 94.5% of the top-1K molecules, greatly outperforming random screening. Importantly, MolPAL retrieves 98.5% and 95.4% of the top-1K given the same model configuration and budget – revealing that factored synthon acquisition degrades performance minimally compared to learning over full molecules for these tasks. Interestingly, SALSA learns faster than MolPAL in early rounds for ROCS-TC, perhaps reflective of the approximate additivity of shape-based scoring across fragments (Cheng & Beroza, 2023). TS (Tab.) is significantly less sample efficient due to its lack of generalization across synthons. The heatmaps in Fig. 2 illustrate SALSA’s progressive ranking of the target space based on the probability of acquiring a given molecule’s constituent synthons. SALSA quickly learns to separate high-scoring and low-scoring molecules, and refines its sampling distribution over multiple rounds to prioritise top molecules, visualised as the top-1K molecules (in red) moving steadily towards the top right corner where acquisition probability is highest.

Scaling beyond enumerable spaces Next, we assess SALSA’s ability to scale to multi-vector spaces beyond the domain of exhaustive screening. We use the same task setup as above, this time subsampling 1K, 10K, and 100K synthons for each vector to construct increasingly large spaces of size 1M, 100M, 10B, and a final $\sim 2T = 910K \times 2.4M$ space, where all synthons are made available. We fix our budget in each space for calls to both the Hybrid Docking and ROCS-TC objective functions to a reasonable 10K molecules per round for 10 rounds of active learning. For completeness, we report runtime for each space in A.9.

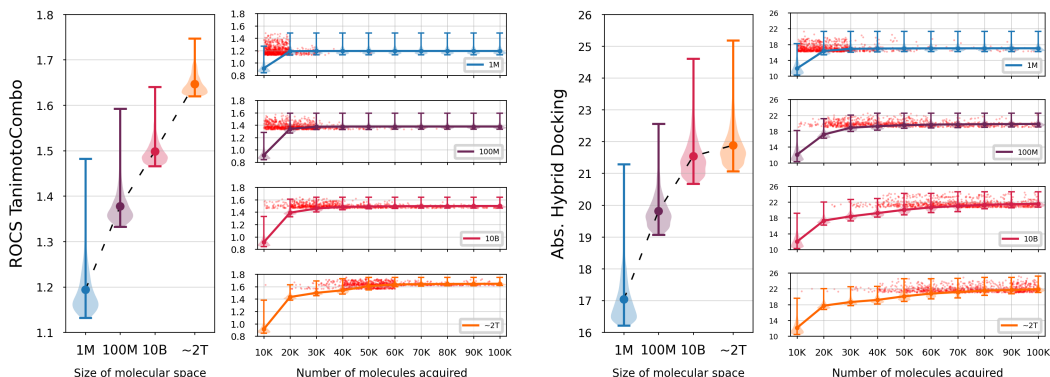


Figure 3: The large violin plots show the min, max, mean, and estimated score density for the top-1K molecules identified by SALSA as space size increases for shape- (left) and structure-based (right) objectives, smoothed over 3 trials. Subplots on the right show the evolution of the top-1K distribution over AL rounds – the final top-1K molecules are marked in red at their sample index.

Fig. 3 demonstrates that SALSA consistently finds better scoring molecules with increasing space size, where we observe an approximately log-linear improvement in ROCS-TC score. The rate of improvement appears to decrease for docking between 10B and $\sim 2T$ molecules. It is informative to look at the acquisition of top-1K molecules (in red) identified during a given run. For ROCS-TC, even the $\sim 2T$ space requires only 60K-70K sampled molecules before diminishing returns. In contrast, many new top molecules appear in the final rounds for docking in both the 10B and $\sim 2T$ spaces. This may indicate either model saturation, or that a 100K learning budget is insufficient to fully explore – consistent with Fig. 2B where learning converges more slowly for docking. However, the upper tail of the top-1K distribution appears to have converged, suggesting the possibility that there are few significantly higher scoring molecules identifiable by SALSA. Lyu et al. (2023) demonstrate a log-linear improvement for top docking scores in increasingly large virtual screens, but in the multi-vector case the core is fixed, potentially constraining the highest achievable score.

Multi-parameter objectives Finally, we explore SALSA’s ability to optimize simple MPOs derived for targets from three protein classes: CDK2, a kinase; BACE1 (PDB: 2IRZ (Rajapakse et al., 2006)), a protease; and DRD2 (PDB: 6LUQ (Fan et al., 2020)), a GPCR. Molecular cores were again extracted from co-crystallized ligands and used to generate synthons (see A.2). We additionally impose light substructure filtering with standard structural alerts to represent a more realistic target space. After filtering, $547K \times 1.3M$, $1.0M \times 1.1M$, and $1.2M \times 246K$ synthons remained for CDK2, BACE1, and DRD2, respectively. Two simple linear MPOs were assessed for each system: Hybrid Docking + QED, and ROCS-TC + QED, with each component scaled to (0, 1), and a 1:1 and 2:1 weighting applied, respectively (see A.7). 10K objective function calls were again budgeted for each of 10 rounds. We plot the MPO components of the top-1K molecules acquired via random acquisition and SALSA in Fig. 4A. We also compare to LibINVENT (Fialkova et al., 2021), an established generative method for multi-vector expansion, using the implementation from the REINVENT4 (Loeffler et al., 2024) framework – allocating 100K objective function calls for parity.

Fig. 4A shows that SALSA achieves equal or better MPO scores to LibINVENT across all targeted tasks. SALSA consistently obtains higher-scoring molecules for both ROCS-TC and docking components compared to LibINVENT. Conversely, LibINVENT produces either similar or slightly improved QED scores. This is likely due to LibINVENT’s prior, which is trained to generate fragments for spliced drug-like ChEMBL molecules, biasing towards QED which is fitted to ChEMBL data (Bickerton et al., 2012). However, Fig. 4B demonstrates that SALSA finds far more unique high-scoring scaffolds compared to LibINVENT across all tasks. We suspect this stems from the use of strict reaction filter penalties to promote synthesizability, sparsifying learning and making it difficult for LibINVENT to move away from its prior. This highlights an important advantage of explicitly optimizing within a targeted synthesizable chemical space. To reinforce this, in A.11 we show that SALSA still outperforms even when applied to spliced ChEMBL fragments.

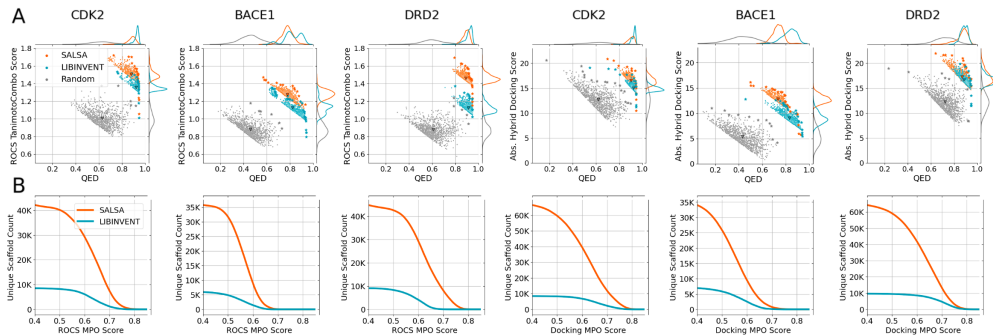


Figure 4: **A.** Top-1K molecules identified across three targets for SALSA, random acquisition, and LibINVENT using QED (Bickerton et al., 2012) plus ROCS-TC and Hybrid Docking objectives. Mean scores and top-20 pareto optimal molecules are denoted by triangles and stars, respectively. **B.** Number of unique Bemis-Murcko scaffolds above a given score value for molecules identified by SALSA and LibINVENT. SALSA compounds show substantially greater diversity.

4 CONCLUSION

SALSA is a sample-efficient and scalable algorithm for virtual screening in non-enumerable multi-vector spaces: our experiments highlight sample efficiency – SALSA identifies approximately 95% of top-1K compounds after evaluating a small fraction of a 1M molecule space for both shape- and ligand-based objectives. Further, SALSA is able to consistently identify increasingly high-scoring molecules for design tasks in spaces up to 2T molecules. SALSA also enables multi-vector screening against MPOs, improving upon the output of a directly comparable generative method in LibINVENT, particularly in terms of diversity. This approach also has qualitative advantages:

Explicit control over chemical space Practitioners can easily inject expert knowledge through explicit and granular control over the target chemical space. Synthon or fragment sets can simply be filtered for desired physico-chemical properties at each vector a priori, as well as for practical considerations such as building block logistics (e.g., cost, lead-time).

Embedded synthetic route Given that each synthon is associated with a reaction, each molecule identified by SALSA comes with a predicted synthesis route which medicinal chemists can evaluate for real-world accessibility. While the enacted synthetic route may evolve in practice, providing these routes as starting points helps streamline the transition from the design to make phase, greatly aiding actionability. Further, the target space can be tuned for stricter practicality by ensuring that the associated reactions are robust and constrained to up-to-date building block catalogues.

Limitations and future directions Successive research in this direction will likely involve iterating on the underlying modeling assumptions and acquisition strategy to more delicately balance exploration and exploitation, continuing to improve sample efficiency. It is also plausible to further increase SALSA’s computational efficiency and scalability by acquiring synthons without exhaustive surrogate model inference while maintaining explicit control over the target space, e.g. by adapting the action space in existing *de novo* methods such as Cretu et al. (2024). Removing this constraint may also enable joint modeling of the synthon space, alleviating the implicit, naive independence assumptions that enable SALSA to scale but risk breaking down against more complex objective functions. We also expect the application of SALSA and similar algorithms to extend to molecular design tasks other than multi-vector expansion, for example to scaffold-hopping, and to screening of ultra-large synthesis-on-demand libraries, such as Enamine’s REAL (Enamine, 2024). We leave these considerations to future work.

REFERENCES

- Louis Bellmann, Patrick Penner, and Matthias Rarey. Topological Similarity Search in Large Combinatorial Fragment Spaces. *Journal of Chemical Information and Modeling*, 61(1):238–251, 2021. ISSN 1549-9596. doi: 10.1021/acs.jcim.0c00850.
- G. Richard Bickerton, Gaia V. Paolini, Jérémy Besnard, Sorel Muresan, and Andrew L. Hopkins. Quantifying the chemical beauty of drugs. *Nature Chemistry*, 4(2):90, 2012. ISSN 1755-4349. doi: 10.1038/nchem.1243.
- John Bradshaw, Brooks Paige, Matt J Kusner, Marwin H S Segler, and José Miguel Hernández-Lobato. A Model to Search for Synthesizable Molecules. *arXiv*, 2019. doi: 10.48550/arxiv.1906.05221.
- Raymond E. Carhart, Dennis H. Smith, and R. Venkataraghavan. Atom pairs as molecular features in structure-activity studies: definition and applications. *J. Chem. Inf. Comput. Sci.*, 25:64–73, 1985. URL <https://api.semanticscholar.org/CorpusID:37017771>.
- Chen Cheng and Paul Beroza. Shape-Aware Synthon Search (SASS) for virtual screening of synthon-based chemical spaces. *ChemRxiv*, 2023. doi: 10.26434/chemrxiv-2023-68jzh.
- Miruna Cretu, Charles Harris, Iliia Igashov, Arne Schneuing, Marwin Segler, Bruno Correia, Julien Roy, Emmanuel Bengio, and Pietro Liò. Synflownet: Design of diverse and novel molecules with synthesis constraints. 2024. URL <https://arxiv.org/abs/2405.01155>.
- Enamine. REAL database, 2024. URL <https://enamine.net/compound-collections/real-compounds/real-database>.
- William Falcon, Jirka Borovec, Adrian Wälichli, Nic Eggert, Justus Schock, Jeremy Jordan, Nicki Skafte, Vadim Bereznyuk, Ethan Harris, Tullie Murrell, et al. Pytorchlightning/pytorch-lightning: 0.7.6 release. *Zenodo*, 2020.
- Luyu Fan, Liang Tan, Zhangcheng Chen, Jianzhong Qi, Fen Nie, Zhipu Luo, Jianjun Cheng, and Sheng Wang. Haloperidol bound D2 dopamine receptor structure inspired the discovery of subtype selective ligands. *Nature Communications*, 11(1):1074, 2020. doi: 10.1038/s41467-020-14884-y.
- Vendy Fialkova, Jiaxi Zhao, Kostas Papadopoulos, Ola Engkvist, Esben Jannik Bjerrum, Thierry Kogej, and Atanas Patronov. LibINVENT: Reaction-based Generative Scaffold Decoration for in Silico Library Design. *Journal of Chemical Information and Modeling*, 2021. ISSN 1549-9596. doi: 10.1021/acs.jcim.1c00469.
- Yarin Gal and Zoubin Ghahramani. Dropout as a bayesian approximation: Representing model uncertainty in deep learning. In Maria-Florina Balcan and Kilian Q. Weinberger (eds.), *Proceedings of the 33rd International Conference on Machine Learning, ICML 2016, New York City, NY, USA, June 19-24, 2016*, volume 48 of *JMLR Workshop and Conference Proceedings*, pp. 1050–1059. JMLR.org, 2016. URL <http://proceedings.mlr.press/v48/gall16.html>.
- Wenhao Gao and Connor W. Coley. The Synthesizability of Molecules Proposed by Generative Models. *Journal of Chemical Information and Modeling*, 60(12):5714–5723, 04 2020. ISSN 1549-9596. doi: 10.1021/acs.jcim.0c00174. URL <https://doi.org/10.1021/acs.jcim.0c00174>.
- Wenhao Gao, Shitong Luo, and Connor W Coley. Generative Artificial Intelligence for Navigating Synthesizable Chemical Space. *arXiv*, 2024.
- David E Graff, Eugene I Shakhnovich, and Connor W Coley. Accelerating high-throughput virtual screening through molecular pool-based active learning. *Chem Sci.*, 2021. ISSN 2021;12(22):7866-7881. doi: 10.1039/d0sc06805e.
- Francesca Grisoni, Berend J. H. Huisman, Alexander L. Button, Michael Moret, Kenneth Atz, Daniel Merk, and Gisbert Schneider. Combining generative artificial intelligence and on-chip synthesis for de novo drug design. *Science Advances*, 7(24):eabg3338, 2021. ISSN 2375-2548. doi: 10.1126/sciadv.abg3338.

- Esther Heid, Kevin P. Greenman, Yunsie Chung, Shih-Cheng Li, David E. Graff, Florence H. Vermeire, Haoyang Wu, William H. Green, and Charles J. McGill. Chemprop: A Machine Learning Package for Chemical Property Prediction. *Journal of Chemical Information and Modeling*, 64(1):9–17, 2024. ISSN 1549-9596. doi: 10.1021/acs.jcim.3c01250.
- Kathryn Klarich, Brian Goldman, Trevor Kramer, Patrick Riley, and W. Patrick Walters. Thompson Sampling: An Efficient Method for Searching Ultralarge Synthesis on Demand Databases. *Journal of Chemical Information and Modeling*, 2024. ISSN 1549-9596. doi: 10.1021/acs.jcim.3c01790.
- Thomas Liphardt and Thomas Sander. Fast Substructure Search in Combinatorial Library Spaces. 2023. doi: 10.26434/chemrxiv-2023-ckbzd.
- Hannes H. Loeffler, Jiazhen He, Alessandro Tibo, Jon Paul Janet, Alexey Voronov, Lewis H. Mervin, and Ola Engkvist. Reinvent 4: Modern AI-driven generative molecule design. *Journal of Cheminformatics*, 16(1):20, 2024. ISSN 1758-2946. doi: 10.1186/s13321-024-00812-5.
- Jiankun Lyu, John J. Irwin, and Brian K. Shoichet. Modeling the expansion of virtual screening libraries. *Nature Chemical Biology*, 19(6):712–718, 2023. ISSN 1552-4450. doi: 10.1038/s41589-022-01234-w.
- Mcule. Mcule Database, accessed 18 Sep 2023. URL <https://mcule.com/database/>.
- David Mendez, Anna Gaulton, A. P. Bento, Jon Chambers, Marleen De Veij, Edwin Félix, María P. Magariños, Juan F. Mosquera, Prudence Mutowo, Michał Nowotka, Marta Gordillo-Marañón, Fiona Hunter, Leandro Junco, Grace Mugumbate, María Rodríguez-Lopez, Francis Atkinson, Nicolas Bosc, Charles J. Radoux, Aldo Segura-Cabrera, Anne Hersey, and Andrew R. Leach. ChEMBL: towards direct deposition of bioassay data. *Nucleic Acids Res*, 47(D1):D930–D940, 2019. doi: 10.1093/nar/gky1075.
- OpenEye. OpenEye Scientific Software, OpenEye toolkits (2022.1.1). <http://www.eyesopen.com>.
- Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, et al. Pytorch: An imperative style, high-performance deep learning library. *Advances in neural information processing systems*, 32, 2019.
- Hemaka A. Rajapakse, Philippe G. Nantermet, Harold G. Selnick, Sanjeev Munshi, Georgia B. McGaughey, Stacey R. Lindsley, Mary Beth Young, Ming-Tain Lai, Amy S. Espeseth, Xiao-Ping Shi, Dennis Colussi, Beth Pietrak, Ming-Chih Crouthamel, Katherine Tugusheva, Qian Huang, Min Xu, Adam J. Simon, Lawrence Kuo, Daria J. Hazuda, Samuel Graham, and Joseph P. Vacca. Discovery of Oxadiazoyl Tertiary Carbinamine Inhibitors of α -Secretase (BACE-1) $\dagger$. *Journal of Medicinal Chemistry*, 49(25):7270–7273, 2006. ISSN 0022-2623. doi: 10.1021/jm061046r.
- RDKit. Open-source cheminformatics. (2023.09.6). <https://www.rdkit.org>.
- Philipp Renz, Dries Van Rompaey, Jörg Kurt Wegner, Sepp Hochreiter, and Günter Klambauer. On failure modes in molecule generation and optimization. *Drug Discovery Today: Technologies*, 32:55–63, 2019. ISSN 1740-6749. doi: 10.1016/j.ddtec.2020.09.003.
- Arman A. Sadybekov, Anastasiia V. Sadybekov, Yongfeng Liu, Christos Iliopoulos-Tsoutsouvas, Xi-Ping Huang, Julie Pickett, Blake Houser, Nilkanth Patel, Ngan K. Tran, Fei Tong, Nikolai Zvonok, Manish K. Jain, Olena Savych, Dmytro S. Radchenko, Spyros P. Nikas, Nicos A. Petasis, Yurii S. Moroz, Bryan L. Roth, Alexandros Makriyannis, and Vsevolod Katritch. Synthon-based ligand discovery in virtual libraries of over 11 billion compounds. *Nature*, 601(7893):452–459, 2022. ISSN 0028-0836. doi: 10.1038/s41586-021-04220-9.
- Robert Schmidt, Raphael Klein, and Matthias Rarey. Maximum Common Substructure Searching in Combinatorial Make-on-Demand Compound Spaces. *Journal of Chemical Information and Modeling*, 2021. ISSN 1549-9596. doi: 10.1021/acs.jcim.1c00640.
- Gisbert Schneider. Automating drug discovery. *Nature Reviews Drug Discovery*, 17(2):97, 02 2017. ISSN 1474-1784. doi: 10.1038/nrd.2017.232. URL <http://www.nature.com/articles/nrd.2017.232>.

Zuofan Wu Jian Peng Connor W. Coley Jianzhu Ma Shitong Luo, Wenhao Gao. Projecting Molecules into Synthesizable Chemical Spaces. 2024. URL <https://arxiv.org/pdf/2406.04628v1>.

Tomer Weiss, Eduardo Mayo Yanes, Sabyasachi Chakraborty, Luca Cosmo, Alex M. Bronstein, and Renana Gershoni-Poranne. Guided diffusion for inverse molecular design. *Nature Computational Science*, 3(10):873–882, 2023. doi: 10.1038/s43588-023-00532-0.

Daniel J. Wood, Svitlana Korolchuk, Natalie J. Tatum, Lan-Zhen Wang, Jane A. Endicott, Martin E.M. Noble, and Mathew P. Martin. Differences in the Conformational Energy Landscape of CDK1 and CDK2 Suggest a Mechanism for Achieving Selective CDK Inhibition. *Cell Chemical Biology*, 26(1):121–130.e5, 2019. ISSN 2451-9456. doi: 10.1016/j.chembiol.2018.10.015.

A APPENDIX

A.1 SALSA ALGORITHM

In algorithm 1, we present pseudo-code for SALSA as implemented in the above experiments, i.e. applied to 2-vector enumeration, with a surrogate model per vector.

Algorithm 1: Scalable Active Learning via Synthon Acquisition

Input: Synthon sets $\mathcal{S}_i$ for $i \in \{0, 1\}$; objective function $f : \text{mols} \rightarrow \mathbb{R}$

Config: Surrogate models $\hat{f}_i(s) \rightarrow \mathcal{N}(\mu(s), \sigma(s))$ for $s \in \mathcal{S}_i, i \in \{0, 1\}$;
 $N \in \text{int}$ rounds; $K \in \text{int}$ samples per round; $\rho_{\max} \in \text{int}$ max sample attempts;
 acquisition strategy $\alpha_f : \mathcal{S} \rightsquigarrow \mathbb{R}$ (i.e. stochastic e.g. for TS $\alpha_f(s) \leftarrow x \sim f(s)$)

Output: $\mathcal{M}_f \subset \text{mols} \times \mathbb{R}$ a set of scored molecules w.r.t. f

Randomly sample K molecules

$\mathcal{M}_{\text{new}} \leftarrow \{(\text{mol}(s_0, s_1) \text{ for } (s_0, s_1) \in \text{zip}(\text{random}(\mathcal{S}_0), \text{random}(\mathcal{S}_1)))\}$

$\mathcal{M}_f \leftarrow \emptyset, \rho \leftarrow 0$ *# Initialize scored set, and sample attempt count*

for $n \leftarrow 1$ **to** N **do**

$\mathcal{M}_f \leftarrow \mathcal{M}_f \cup \{(m, f(m)) \text{ for } m \in \mathcal{M}_{\text{new}}\}$ *# Score new molecules*

$\mathcal{M}_{\text{new}} \leftarrow \emptyset$

if $n < N$ **and** $\rho \leq \rho_{\max}$ **then**

$\mathcal{D}_i \leftarrow \{(m_i, y) \text{ for } (m, y) \in \mathcal{M}_f\} \text{ for } i \in \{0, 1\}$ *# Update synthon datasets*

$f_i^* \leftarrow \hat{f}_i \cdot \text{fit}(\mathcal{D}_i) \text{ for } i \in \{0, 1\}$ *# Train surrogate models*

Acquire new molecules via synthons

ρ will terminate loop if new molecules are sampled too infrequently (i.e. convergence)

$\rho \leftarrow 0$

while $\text{len}(\mathcal{M}_{\text{new}}) < K$ **and** $\text{count} < \rho$ **do**

$s_i^* \leftarrow \text{argmax}_{s \in \mathcal{S}_i} \alpha_{f_i^*}(s) \text{ for } i \in \{0, 1\}$

$\rho \leftarrow \rho + 1$

if $\text{mol}(s_0, s_1) \notin \mathcal{M}_{\text{new}} \cup \{m \text{ for } (m, -) \in \mathcal{M}_f\}$ **then**

$\mathcal{M}_{\text{new}} \cdot \text{add}(\text{mol}(s_0, s_1))$ *# Only add sampled molecule if unseen*

return $\mathcal{M}_f$

Where: $\text{mol} : \mathcal{S}_0 \times \mathcal{S}_1 \rightarrow \text{mols}$ and for convenience $\text{mol}(s_0, s_1)_i := s_i$

A.2 SYNTHON SPACES

Molecular scaffolds In our experiments, SALSA was applied in the context of multi-vector enumeration from an explicit core. We extracted cores from ligands found in pertinent PDB structures containing at least one ring system with two R-groups for optimization, specifically from 3D co-crystallised systems in order to facilitate shape- and structure-based scoring. We identified CDK2 (PDB: 6GUH)(Wood et al., 2019), BACE1 (PDB: 2IRZ) (Rajapakse et al., 2006), and DRD2 (PDB: 6LUQ) (Fan et al., 2020) as suitable candidates to represent design tasks. Each extracted core was replaced with a synthetic intermediate with functionalized reaction handles at the target R-groups to enable synthon generation for our experimental target spaces (see Fig. 5).

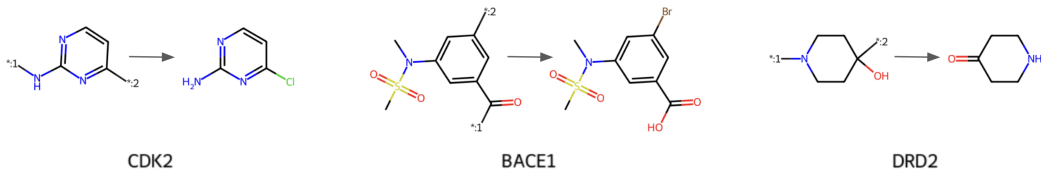


Figure 5: Core scaffolds mapped to synthetic intermediates with functionalized reaction handles.

Synthon space construction For our experiments, we explicitly construct large multi-vector spaces. Given a core scaffold with R-group handles at desired vectors, we determine applicable reactions via partial substructure matching over a set of bimolecular SMIRKS reactions. Building blocks that match the corresponding pattern are retrieved from a database of commercially available options via Molecule (accessed 18 Sep 2023). Synthons are created by replacing displaced reacting groups with a generic linker atom, as detailed in Liphardt & Sander (2023).

A.3 MPNN HYPERPARAMETERS

We minimally adapt the default `chemprop` (Heid et al., 2024) architecture: a message-passing depth of 3 for the encoder, 2 layers for the MVE head, and 300 hidden dimensions with ReLU activations throughout for both. In each round, we trained from scratch for a maximum of 50 epochs with a batch size of 64, holding out 20% data for early-stopping on the MVE validation loss with patience=10. We optimized with Adam and a NoamLR scheduler with initial, max, and final LR of $1e-4$, $1e-3$, $1e-3$, respectively. The MPNN itself is implemented in PyTorch (Paszke et al., 2019) and trained via lightning (Falcon et al., 2020).

A.4 ALTERNATIVE ACQUISITION STRATEGIES AND SURROGATE MODELS

In Fig. 6 we ablate MolPAL’s implementation of random forest (RF) and feed-forward neural network (NN) models against the MPNN model used in our experiments, using 2048-bit atom-pair fingerprints with a minimal and maximal radius of 1 and 3 (Carhart et al., 1985) as fixed features. Here, the acquisition strategy is held fixed to ϵ -greedy, where top synthons are chosen except for an $\epsilon = 5\%$ chance of picking at random. We see that the MPNN significantly outperforms the NN and RF. We also report performance of a number of non-stochastic acquisition strategies used in MolPAL, including probability of improvement (PI), expected improvement (EI), upper confidence bounds (UCB), and a “non-stochastic” TS where scores are drawn only once (see Graff et al. (2021) for definitions). Throughout, a molecule’s acquisition score was defined as the sum of its synthon acquisition scores. We found that stochastic TS performed best. This is consistent with the promising performance of the method in Klarich et al. (2024), which is conceptually similar to this configuration of SALSA, instead using tabulated predictions with online TS. This suggests that optimistic, exploratory strategies are most effective when navigating via synthon spaces in this manner.

We also trialed inference-time dropout as an alternative to MVE for uncertainty quantification, training to predict the mean via MSE loss, and estimating the mean and variance by sampling 10 predictions with a dropout probability of 0.2. This method can be interpreted as an approximate Bayesian (i.e. variational) technique for modeling epistemic uncertainty (Gal & Ghahramani, 2016). We believe the drastically inferior performance of this approach when compared to its application in Graff et al. (2021) stems from the aleatoric variance of a given synthon’s score distribution dominating the epistemic variance in its predicted mean, due to the unobserved choice of complementary synthon.

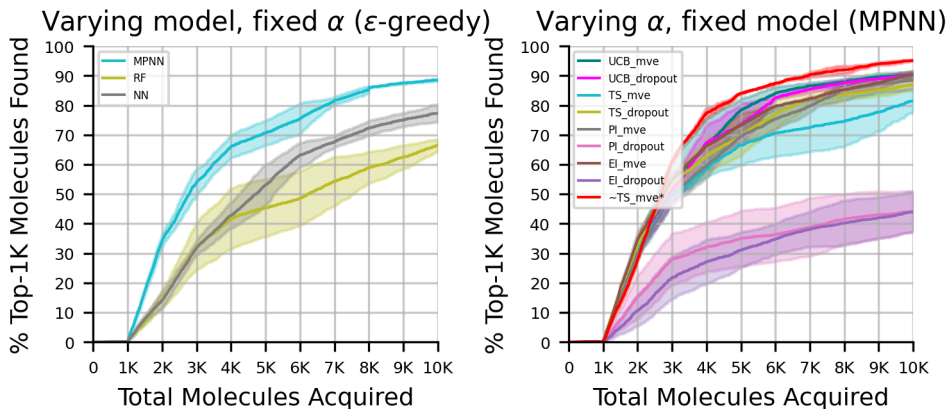


Figure 6: Ablation Acquisition of the top-1K scoring compounds on the 1M-space ROCS-TC design task. Each method was allocated a 1K objective budget per round for 10 rounds.

A.5 SALSA WITH ONE MODEL

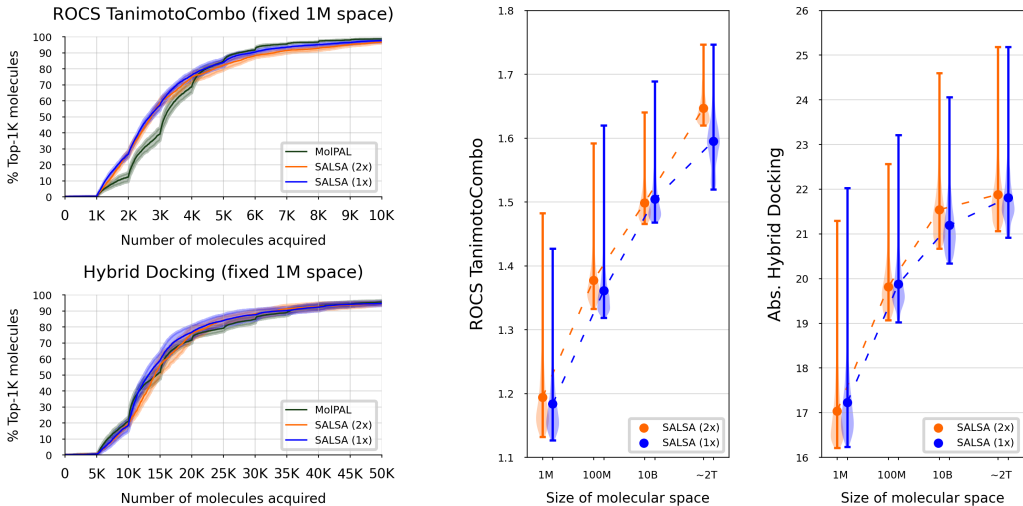


Figure 7: One vs two-model SALSA for sample efficiency (5 trials) and scaling (3 trials)

In Fig. 7, we repeat our sample efficiency and scaling experiments from 3 using a single model configuration. Here, we simply add a one-hot encoding to the learned graph embedding to indicate the vector to which each synthon belongs, and the model sees batches of data from both vectors during training. We did not modify the model architecture or training hyperparameters. We find that the performance of the single-model variant is similar or slightly improved in the benchmark space, ultimately recalling a mean of 97.5% and 94.6% of the top-1K molecules, compared to 96.5% and 94.5% for two models with respect to ROCS-TC and Hybrid Docking (MolPAL: 98.5% and 95.4%). The scaling experiments show comparable trends for both variants as the size of the molecular space increases, but SALSA (1x) appears to perform slightly worse in the largest space for ROCS-TC. We hypothesize that the single model variant may saturate faster, requiring greater capacity to handle its larger data distribution. We highlight the possibility of using a single model due to its advantages: simpler, more efficient use of computational resources, and transferability to other molecular design contexts such as de novo design, which may additionally benefit from more informative featurization and generalization across different synthon sets.

A.6 BASELINES

MolPAL In our sample efficiency experiments, we run MolPAL in its best reported configuration in Graff et al. (2021), with an MPNN surrogate model configured as in A.3, and greedy acquisition.

(Tabular) Thompson sampling We modified the open source implementation of Klarich et al. (2024) at <https://github.com/PatWalters/TS>, adding a custom CSV evaluator class to enable scoring with pre-computed scores. We provided the same synthon sets used by SALSA as the reagent lists. We allowed for 3 and 2 warm-up trials for Hybrid Docking and ROCS-TC scoring, respectively. We did not count the warm-up trials towards the objective function budget.

LibINVENT In our MPO experiments, we configured the LibINVENT implementation in Loeffler et al. (2024) in staged learning mode, running for 1600 iterations (`max_steps`) with a `batch_size` of 64 to generate 100K molecules. We used default $\sigma=128$ for the Difference between Augmented and Posterior (DAP) reward strategy, and a learning rate of $1e-4$.

A.7 MOLECULAR SCORING FUNCTIONS

Docking: OpenEye Hybrid Docking Prior to docking, sampled molecules undergo preparation, including protomer, stereochemistry, and conformer generation using `openeye-toolkits` (version 2022.1.1) (OpenEye). Protomer generation was performed using OpenEye Quacpac. Stereochemistry enumeration and conformer generation were carried out using OpenEye Omega, with a maximum of 10 stereocenters and 200 conformers per molecule. Following molecular preparation, docking was conducted using OpenEye Hybrid. For CDK2, BACE1, and DRD2, the outputted scores were divided by a factor of -24, -17, and -24, to scale roughly between 0 and 1. The scaled docking score of the highest scoring conformer was assigned to the molecule.

ROCS-TC: OpenEye TanimotoCombo Score For ROCS-TC, all molecules undergo preparation similarly to docking. We use OpenEye ROCS to perform shape similarity scoring. We again define a molecule’s score to be the highest TanimotoCombo Score achieved by any of its conformers, where TanimotoCombo Score consists of an equally weighted sum of shape and colour Tanimoto scores. We divided these scores by 2 during learning, to scale the output range between 0 and 1.

QED QED (Bickerton et al., 2012) was calculated using RDKit (RDKit). Each enumerated molecule was converted into an RDKit molecule and scored using the RDKit QED function.

MPO objective functions We defined two simple linear MPOs for each of the three protein targets – CDK2, BACE1, and DRD2. ROCS-TC + QED was weighted 2:1. In Docking + QED, QED was weighted equally to the scaled docking scores.

A.8 ChEMBL MOLECULES, PHYSIO-CHEMICAL CALCULATIONS AND ADMET PREDICTIONS

In Fig. 8, the top-1K molecules from SALSA and LibINVENT runs are compared to the top-1K bioactive molecules in ChEMBL as measured by pChEMBL score for each protein target. Properties were calculated via RDKit where possible, or otherwise predicted (via internal models) in order to assess drug-likeness over twelve metrics: molecular weight (MW), total polar surface area (TPSA), hydrogen bond acceptors (HBA), hydrogen bond donors (HBD), number of aromatic rings (AROM), number of structural alerts (ALERTS), predicted lipophilicity at pH 7.4 (LogD), number of rotatable bonds (ROTB), hERG potency (hERG), PXR potency (PXR), log fraction of ligand unbound in human plasma (Fraction Unbound in Plasma), and CACO2 permeability (CACO2). We observe the majority of molecules generated by both methods fall within desirable drug-like space and within similar or better bounds than molecules retrieved from ChEMBL.

ChEMBL Bioactive molecules associated with each protein target were downloaded from the ChEMBL33 database (Mendez et al., 2019). Molecules without associated SMILES values were removed. The top-1K (unique) molecules with highest pChEMBL values for each target were selected, where pChEMBL value is the negative logarithm of the molar IC50, EC50, Ki, Kd, or Potency.

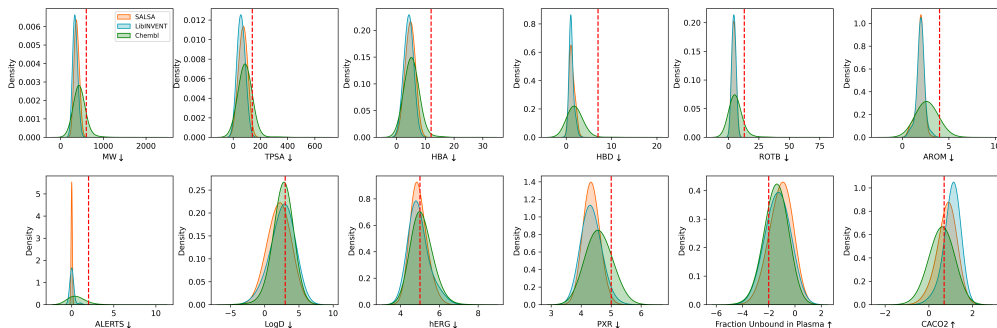


Figure 8: ADMET properties for the top-1K molecules generated by SALSA and LibINVENT compared with the top-1K bioactive molecules from ChEMBL for each protein target, ranked by pChEMBL value. Dashed lines represent thresholds for these properties, and an up or down arrow represents preference for values greater or less than the threshold, respectively. Density estimates are aggregated over all MPOs and protein targets to give a high-level view on property distributions.

A.9 RUNTIME AND MODEL INFERENCE

SALSA is able to effectively screen large multi-vector spaces within a matter of hours using a single A10G GPU for training and inference. Time spent on training, scoring, and acquisition is primarily a function of the number of total molecules acquired and the number of rounds. Model inference scales linearly with the size of the targeted synthon/fragment sets. Time spent on scoring, model training, inference, and acquisition for different space sizes in our scaling experiments in 3 is outlined in Tables 1 and 2. Here, inference time remains negligible until the $\sim 2T$ space, where it begins to take ~ 2 hours to compute the necessary statistics for each synthon. This is significant, but it is a million-fold improvement over an equivalent full-molecular pool-based screen requiring one forward pass per molecule. Inference times can easily be improved by distributed compute.

Table 1: ROCS-TC Results

	1M space	100M space	10B space	2T space
scoring	2H 28m 41s $\pm$ 6m 25s	1H 41m 43s $\pm$ 6m 24s	1H 40m 15s $\pm$ 18m 36s	1H 23m 59s $\pm$ 10m 31s
training	2H 21m 37s $\pm$ 22m 22s	2H 13m 42s $\pm$ 3m 48s	2H 15m 38s $\pm$ 4m 9s	2H 6m 15s $\pm$ 1m 20s
inference	0H 0m 11s $\pm$ 0m 0s	0H 0m 59s $\pm$ 0m 1s	0H 8m 36s $\pm$ 0m 15s	2H 15m 46s $\pm$ 3m 52s
overall	5H 11m 43s $\pm$ 20m 16s	4H 5m 0s $\pm$ 4m 36s	4H 20m 55s $\pm$ 20m 30s	8H 28m 7s $\pm$ 9m 31s

Table 2: Hybrid Docking Results

	1M space	100M space	10B space	2T space
scoring	4H 3m 14s $\pm$ 18m 15s	3H 25m 49s $\pm$ 12m 32s	3H 41m 19s $\pm$ 5m 10s	3H 24m 14s $\pm$ 11m 59s
training	2H 46m 46s $\pm$ 4m 43s	2H 36m 32s $\pm$ 9m 10s	2H 41m 24s $\pm$ 9m 1s	2H 37m 13s $\pm$ 6m 42s
inference	0H 0m 13s $\pm$ 0m 1s	0H 0m 57s $\pm$ 0m 1s	0H 8m 45s $\pm$ 0m 11s	2H 21m 0s $\pm$ 4m 12s
overall	7H 0m 9s $\pm$ 17m 9s	6H 12m 12s $\pm$ 8m 55s	6H 47m 52s $\pm$ 5m 2s	10H 45m 16s $\pm$ 4m 58s

A.10 TOP SCORING MOLECULES VISUALIZED

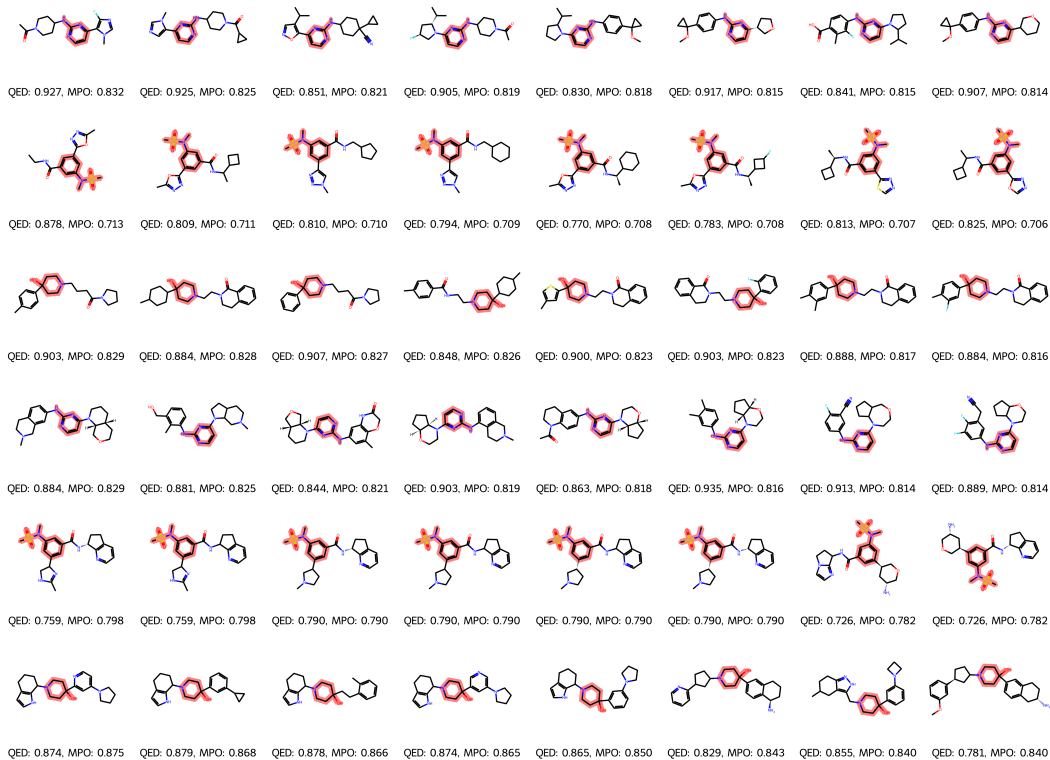


Figure 9: Selected top molecules from SALSA MPO runs for all three protein targets. Highlighted atoms and bonds represent fixed cores. QED and MPO score are labeled below each molecule.

A.11 MPOS IN ChEMBL FRAGMENT SPACE

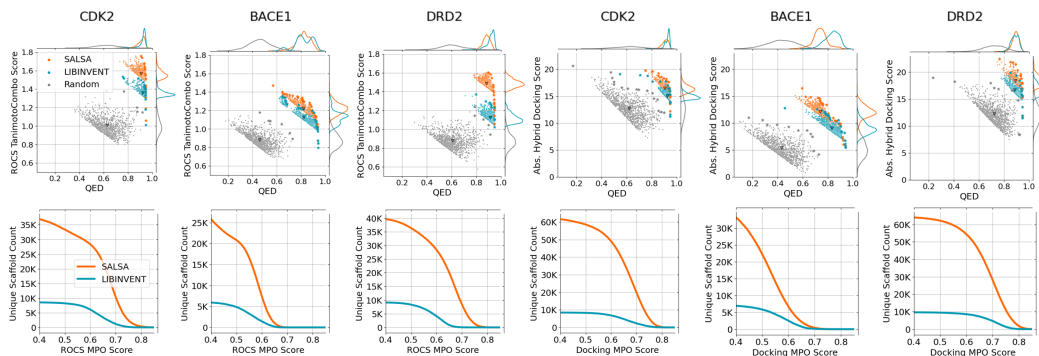


Figure 10: Top: the top-1K molecules identified across three targets for SALSA and LibINVENT using QED (Bickerton et al., 2012) plus ROCs-TC and Hybrid Docking objectives. Bottom: the number of unique scaffolds scoring above a given score for molecules enumerated by SALSA and LibINVENT using ROCs/QED MPO (left) and Docking/QED MPO (right) for each protein target. These runs were performed on fragment spaces formulated from ChEMBL which were generated by breaking all acyclic bonds, and then filtering for ≤ 20 heavy atoms, ≤ 4 H-bond donors, ≤ 4 H-bond acceptors, ≤ 2 rotatable bonds, ≤ 1 ring system, and an observed frequency count of ≥ 10 .