

Deep Batch Active Learning for Drug Discovery

Reviewed Preprint

Published from the original preprint after peer review and assessment by eLife.

[About eLife's process](#)

Reviewed preprint posted

September 13, 2023 (this version)

Posted to bioRxiv

July 29, 2023

Sent for peer review

June 13, 2023

Michael Bailey , Saeed Moayedpour, Ruijiang Li, Alejandro Corrochano-Navarro, Alexander Kötter, Lorenzo Kogler-Anele, Saleh Riahi, Christoph Grebner, Gerhard Hessler, Hans Matter, Marc Bianciotto, Pablo Mas, Ziv Bar-Joseph , Sven Jager 

R&D Data & Computational Science, Sanofi, Cambridge, MA, United States • Digital Data, Sanofi, Shanghai, China • Synthetic Molecular Design, Integrated Drug Discovery, Sanofi-Aventis Deutschland GmbH, Industriepark Höchst, Building G838, 65926 Frankfurt am Main, Germany • Molecular Design Sciences, Integrated Drug Discovery, Sanofi, Vitry-sur-Seine, 94403, France

 https://en.wikipedia.org/wiki/Open_access

 <https://creativecommons.org/licenses/by/4.0/>

Abstract

A key challenge in drug discovery is to optimize, in silico, various absorption and affinity properties of small molecules. One strategy that was proposed for such optimization process is active learning. In active learning molecules are selected for testing based on their likelihood of improving model performance. To enable the use of active learning with advanced neural network models we developed two novel active learning batch selection methods. These methods were tested on several public datasets for different optimization goals and with different sizes. We have also curated new affinity datasets that provide chronological information on state-of-the-art experimental strategy. As we show, for all datasets the new active learning methods greatly improved on existing and current batch selection methods leading to significant potential saving in the number of experiments needed to reach the same model performance. Our methods are general and can be used with any package including the popular DeepChem library.

eLife assessment

This **valuable** study reports a new method based on batch active learning to optimize the biological and pharmaceutical properties of small molecules of pharmaceutical interest. The new method seems **compelling**, but the theoretical analysis is **incomplete** and the reproducibility and impact of the article would benefit from disclosing the code and datasets used in the study. With these aspects strengthened, this paper would be of interest to computational and medicinal chemists and scientists working in the drug discovery field.

The process that leads from a small molecule that shows some activity against the target of interest to a candidate for clinical development involves complex multi-parameter optimization. In this process, in addition to the activity against the target itself, the ADMET profile of the molecule, i.e. its Absorption, Distribution, Metabolism, Excretion, and Toxicity properties are optimized. Accurate *in silico* models for the desired properties are required to speed up and improve decision making and reduce the number of necessary experiments [Hessler and Baringhaus \(2018\)](#); [Grebner et al. \(2022\)](#); [Wu et al. \(2020\)](#). Amongst other machine learning (ML) techniques, deep learning models, and more specifically (graph) neural networks, have been used successfully in this field [Xiong et al. \(2021\)](#). Scale-effective ADMET and affinity prediction methods require an abundance of labeled training data due to their complexity and the need to cover the enormous molecular design space. This can be a bottleneck in cost, time, and experimental resources [Chuang et al. \(2020\)](#). Various approaches have been proposed for improving data acquisition and the models, including transfer learning [Weiss et al. \(2016\)](#), data augmentation [Cai et al. \(2020\)](#) and active learning [Cohn et al. \(1996a\)](#). Current optimization approaches often work in cycles. In each cycle a set of molecules are tested, the model is revised and, based on the revised model, a new set is selected for testing, and so on until the model reaches the desired performance.

Active learning [Cohn et al. \(1996b\)](#) is an approach for selecting molecules for each of the cycles. Unlike traditional approaches that test the most promising candidates in each round [Cohn et al. \(1994\)](#) in active learning samples are selected to optimize the *model* rather than the cycle result. In such approach optimization samples are prioritized by their ability to improve model performance when labeled. Active learning has been studied in sequential mode, where samples are labelled one-at-a-time, and batch mode, where samples are selected for labelling in batches [Settles \(2012\)](#). Batch mode is both more realistic for small molecule optimization and more challenging computationally. The main problem is that samples are not independent (sharing chemical properties that will influence model parameters) and so selecting a set based on marginal improvement does not reflect well the improvement provided by the entire batch [Ash et al. \(2021\)](#).

Active learning methods have been utilized to predict and optimize the physiochemical and biological properties of molecular systems. For instance, batch active learning (BMDAL) was used to enhance the model accuracy in predicting the conformations, energetics, and interatomic forces of small organic compounds where the diversity of the training data is a limiting factor [Zaverkin et al. \(2022\)](#). Active learning was also combined with pairwise difference regression (PADRE) [Tynes et al. \(2021\)](#) to predict molecular properties including redox free energy. Reker et al developed an active learning based tool to screen and select top candidates for ligand-target binding prediction [Reker et al. \(2017\)](#). Thompson et al. [Thompson et al. \(2022\)](#) employed Active learning framework to build a package for the binding free energy of to TYK2 Kinase. In another study model uncertainty for unlabeled pharmacokinetics was used to set an active learning pipeline for plasma exposure prediction [Ding et al. \(2021\)](#). Pertusi et al [Pertusi et al. \(2017\)](#), used an active learning data selection method for characterizing enzyme promiscuity. However, while some of these methods worked well, they were not applied to the more advanced deep learning methods that have become the tool of choice for small molecule modeling.

On the theoretical side, a number of Batch Active learning methods have recently been developed though these have not been used in the drug design space. For example, BAIT [Ash et al. \(2021\)](#) uses a probabilistic approach for the learning procedure that optimally selects (using greedy approximation) a set of samples that maximizes the likelihood of the model parameters (last layer) as defined by Fisher information [Ash et al. \(2021\)](#). Others have proposed using the local approximation to estimate the maximum of the posterior distribution over the batch. This is achieved through computation of the inverse Hessian of the negative log posterior [Daxberger et al.](#)

(2021 [↗](#)). Recently, GeneDisco [Mehrpour et al. \(2021 \[↗\]\(#\)\)](#) was published as an open source library of benchmarking data pertaining to transcriptomics active learning work. However, these methods have not been extensively tested for small molecules optimization. Specifically, to date, popular *in silico* design suits, including ChemML [Haghighatlari et al. \(2019 \[↗\]\(#\)\)](#) and DeepChem [de \(2016\)](#), do not support active learning strategies.

To address these shortcomings, we developed a new innovative and generalizable strategy that can be used with any deep learning ADMET methods. Our methods are inspired by the Bayesian deep regression paradigm, where estimating the model uncertainty is tantamount to obtaining the posterior distribution of the model parameters [Kendall and Gal \(2017 \[↗\]\(#\)\)](#). Model uncertainty is determined using innovative sampling strategies, with no extra model training required. We next select batches that maximize the joint entropy, i.e., the log-determinant of the epistemic covariance of the batch predictions. This enforces batch diversity by rejecting highly correlated batches.

To evaluate our methods and compare them to state of the art approaches we have assembled a large collection of datasets. As we show, our active learning algorithm consistently leads to the best performance when compared to prior methods suggested for this task. We also generated new data for internal candidates and show that our methods can significantly save cost and time compared to current industry optimization approaches for these datasets.

Results

We developed and tested several batch active learning selection approaches which quantify the uncertainty over multiple samples. Given these uncertainties our method aims to select the subset of samples with maximal joint entropy (i.e., information content) (**Figure 1 [↗](#)**). Specifically, we use multiple methods to compute a covariance matrix, C , between predictions on unlabeled samples, $\mathcal{U}$. Then, using an iterative and greedy approach, the method selects a submatrix C_B of size $B \times B$ from C with maximal determinant. Such approach takes into account both the “uncertainty” (which is manifested in the variance of each sample) and the “diversity” (which is reflected in the covariance). See Methods for details.

Evaluating active learning method on ADMET and affinity related data

We used several public drug design datasets to test and compare the performance of our methods. For comparison, in addition to the two methods we developed (MC dropout and Laplace Approximation, COVDROP and COVLAP, respectively) we also compared to methods that have been previously suggested (*k*-means [Nguyen and Smeulders \(2004 \[↗\]\(#\)\)](#) and BAIT [Ash et al. \(2021 \[↗\]\(#\)\)](#)) and to a random ordering of the experiments (i.e. no active learning). Batch size was set to 30 for all methods. During each iteration of the loop, each model (e.g., *k*-Means, BAIT, Random, COVDROP, or COV-LAP) selects a batch consisting of a fixed number of samples from the unlabeled pool. This iterative process is repeated until all labels in the Oracle are exhausted. In the retrospective setting, the pool includes all samples from the relevant dataset, while the oracle retains the corresponding labels. Our evaluation datasets included a cell permeability dataset with 906 drugs [Wang et al. \(2016 \[↗\]\(#\)\)](#), aqueous solubility dataset, which comprises 9,982 small molecules [Sorkun et al. \(2019 \[↗\]\(#\)\)](#); [Huang et al. \(2022 \[↗\]\(#\)\)](#), and the lipophilicity data for 1200 small molecules [Wenlock and Tomkinson \(2015 \[↗\]\(#\)\)](#); [Wu et al. \(2018 \[↗\]\(#\)\)](#). We also included 10 large affinity datasets, 6 from ChEMBL and 4 new internal datasets. Details for all datasets are provided in Methods and in **Table 1 [↗](#)**.

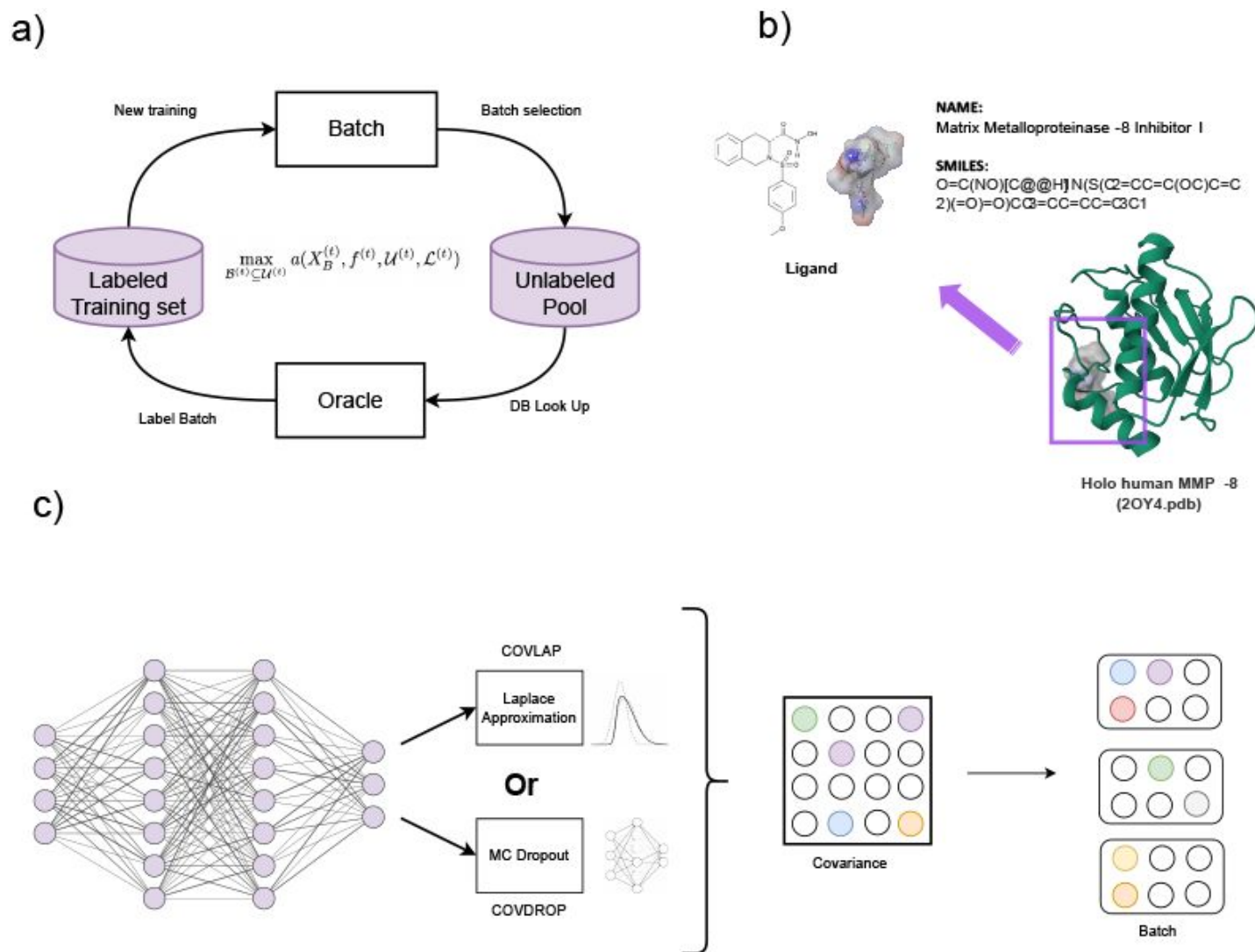


Figure 1.

Method overview. a) The active learning loop. We consider the entire dataset as the pool, and the oracle as the entity that holds the corresponding labels. During each iteration of the active learning loop, we select a batch of points from the pool and obtain their corresponding labels from the oracle. We then update our model using the selected batch and evaluate its performance on a holdout set. This process is repeated until the desired level of performance is achieved. b) Prediction of binding affinity is the target function for the ChEMBL and Sanofi-Aventis datasets c) Active learning batch selection. At the last layer of our model, we use either Laplace approximation or Monte Carlo dropout to compute covariances (COVLAP and COVDROP), from which an ensemble of predictions is generated. With the derived covariance matrix, we optimize batches iteratively based on their information content.

Dataset	COVDROP	COVLAP	<i>k</i> -means	BAIT	Random	Chron.	NC	% gain
Caco2	420	450	570	540	600	-	637	30.0
Lipophilicity	1440	1500	1560	1650	1650	-	2940	12.7
HFE	360	450	450	450	450	-	450	20.0
Solubility	2160	6810	5850	6988	5730	-	6988	62.3
PPBR	1050	450	450	420	510	-	1130	-
GSK3 β	1260	2325	2325	2325	2325	-	2325	45.8
MMP3	690	810	900	870	930	-	1280	25.8
PPAR δ	570	660	660	630	690	-	879	17.4
A1AR	570	540	600	630	630	-	952	9.5
NAV17	1770	4004	4004	4004	4004	-	4004	55.8
D2	2010	4797	4797	4797	4797	-	4797	58.1
Renin	150	240	270	300	330	390	422	54.5
FXa	750	1020	870	1026	990	1026	1026	24.2
MMP-8	150	270	390	420	456	360	456	67.1
PPAR δ	420	300	570	390	570	600	614	26.3

Table 1.

Number or required experiments for reaching model error with 10% higher than the minimum calculated RMSE for the whole training set across different methods. NC - Number of compounds in this dataset. Chron - analysis performed with batch selection using the actual order in which the data was profiled (only available for internal data). % gain estimates the improvement over the Random selection computed from $(1 - N_{\text{COVDROP}}/N_{\text{Random}}) \times 100$.

Comparison of active learning methods for solubility datasets

Results for a selected subset of the benchmarked datasets are presented in [Figure 2](#). In these figures we present the accuracy of models when using the different active learning methods as a function of the iteration. As can be seen, in most cases the COVDROP method very quickly leads to better performance when compared to other methods.

The overall shape of the RMSE profiles is impacted by the statistics of the target values in each dataset. For instance, for the plasma protein binding rate (PPBR) dataset, one can observe that all methods are suffering from large RMSE values early on. Specific to this case, there is an extreme imbalanced distribution for the target value of the source, as illustrated in [Figure 4](#). Using a small number of compounds, the model gets a good insight of the most representative range, with a small peak following within the 300-400 samples, indicating a lack of training in underrepresented regions in early stages. In contrast to PPBR, hydration free energy (HFE) and effective cell permeability (Caco-2) the target distributions suffer less from skewness and spreadiness, as shown in [Figure 4](#). For HFE and Caco-2 datasets COVDROP is the clear winner reaching RMSE within 10% of the final RMSE after testing only 36 and 450 compounds, respectively.

We also tested much larger datasets. For example, the aqueous solubility dataset (AS) ([Figure 1a](#)) is significantly larger than those presented in [Figure 2](#). There we see much slower convergence of RMSE values which we attribute to the normal distribution of the target values ([Figure 4](#)) and larger size of the training data. This RMSE profile indicates that for the Solubility dataset, the BAIT method exhibits inferior performance compared to the other methods, while COVDROP demonstrates the smallest root mean square error (RMSE) on the full dataset compared to all other methods. Additionally, COVDROP outperforms all other methods starting at 400 compounds.

Performance on affinity data

We next evaluated the methods on several affinity datasets from ChEMBL and Sanofi-Aventis with different protein targets. For each of these, a diverse set of ligands is screened for their affinity to a target protein, e.g. MMP-8. [Figure 2c](#) depicts the retrospective experiment for modeling the affinity to Glycogen synthase kinase-3 β (GSK3 β), and we observe a similar trend to the solubility experiments. COVDROP outperforms very early the other methods.

[Figure 2d](#) presents to the retrospective experiment predicting the activity of small molecules against Matrix-metalloprotease(ChEMBL) with COVDROP gain the best option.

The internal datasets provide the opportunity to compare batch selection methods with the actual order used to experimentally test the compounds. We observe that all batch selection methods (though not random) outperform the current human based ordering ([Figure 2 e-f](#) and [1 h-i](#)). For instance, for Renin and MMP-8, COVDROP significantly outperforms the chronological batch selection by requiring 62 and 58 % less training data points.

Similar to the ChEMBL data, we observed that for most of the molecules tested within the Sanofi-Aventis datasets COVDROP significantly outperforms other methods. However, for this dataset we also observed very good performance of COVLAP on the FXa dataset ([Figure 2f](#)). Furthermore, the number of experiments to achieve the 10% higher than minimum RMSE threshold is 750, while other selection methods require at least 12% more experiments to obtain the same metric ([Table 1](#)). [Figure 2e](#) showcases the retrospective experiment on the Renin dataset. Similar to our previous observations, we observed that COVDROP significantly outperformed the other methods. Notably, the method exhibited a small root-mean-square error (RMSE) after the initial batch selection in comparison to the final RMSE attained, which was achieved by the other methods after over 200 experiments and multiple iterations.

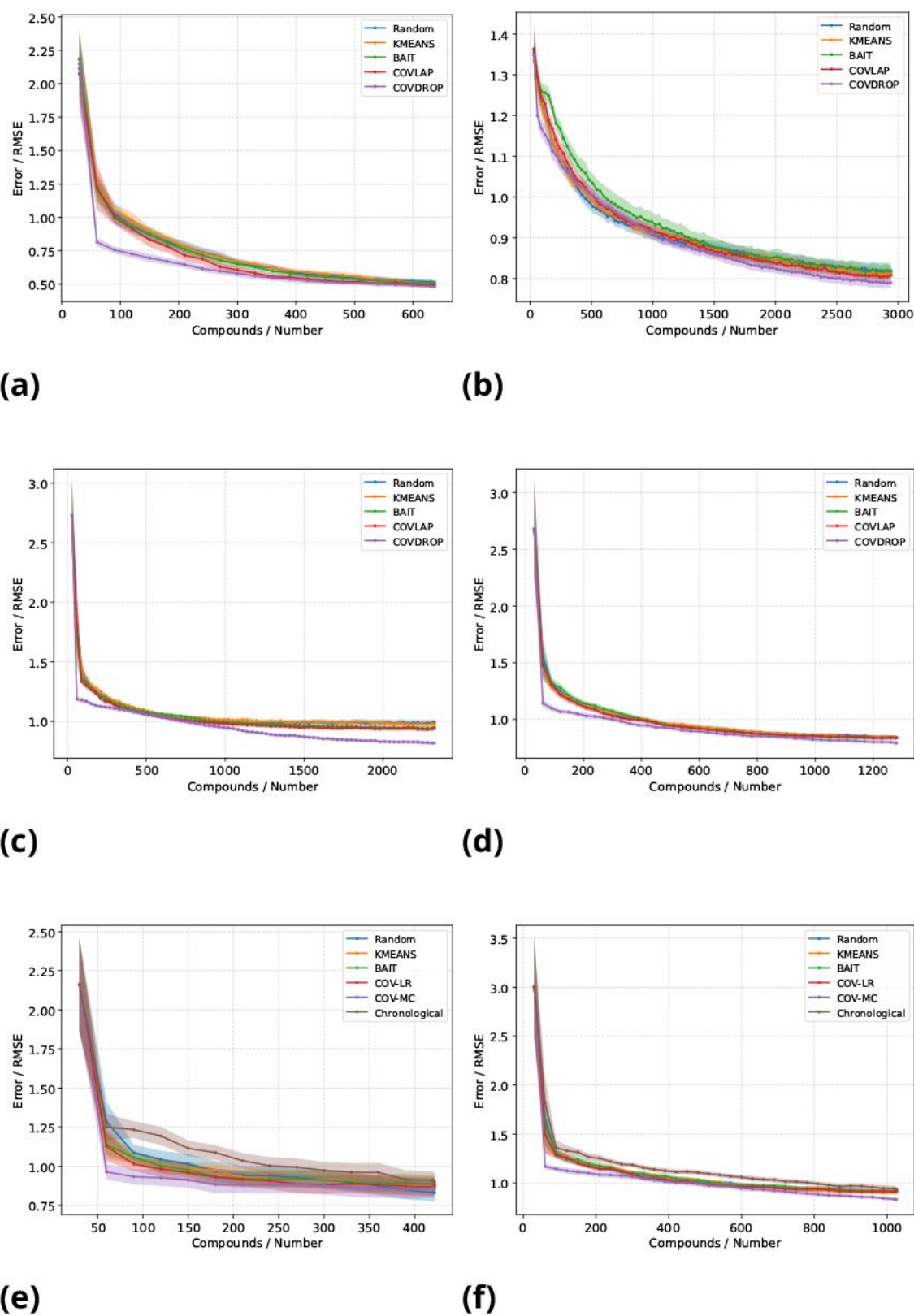


Figure 2.

Performance of different methods for optimizing models for ADMET and affinity related datasets and batch selection strategies: a) the cell effective permeability b) lipophilicity. Panels c-f are affinity measurements to various targets proteins: c) GSK3 β (ChEMBL) d) MMP3 (ChEMBL) e) Renin (Sanofi-Aventis) f) FXa (Sanofi-Aventis)

To further quantify the improvement we present in **Table 1** the number of experiments required by each method for reaching model with an error at most 10% higher than the minimum RMSE obtained using all the data across all selection methods. As the table shows, both of our methods outperform all other methods for most datasets, in some cases significantly so. For example, we observed that for smaller datasets COVDROP can improve the results very quickly leading to much better performance. For most of the datasets performance improvements are greater than 10% vs. Random selection. As mentioned before for PPBR dataset due to the imbalance nature of the data COVDROP underperforms and BAIT selection produces the beset result. Even when compared to the chronological order in which the internal compounds were tested (i.e. to the actual experimental cycle) we observe improvements of up to 62%. This holds true if we change the stopping criteria to 20% or 5% difference (Table 4 & 5).

Methods

Datasets

To benchmark the various batch selection methods, we have collected both private (Sanofi-owned) and public datasets that represent a diverse range of some of the most important molecular properties that scientists need to address when developing small-molecule drugs. We list below the benchmark datasets used in this work covering the properties related to the absorption and distribution pharmacokinetic processes. In addition to the ADMET related properties, the benchmark datasets include four Sanofi and six public datasets recording the affinity of small drug molecules to ten target proteins, such as kinase and GPCRs. **Table 1** provides detailed information on all of these datasets.

Affinity datasets

Affinity measures the strength of binding between the ligands and biological targets, such as proteins. It is a critical molecular property that determines the drug efficacy. There-fore, to validate our active learning strategy for building statistical models, several datasets from the public database ChEMBL [Gaulton et al. \(2012\)](#) and internal sources were used.

ChEMBL datasets

To retrieve suitable datasets from ChEMBL (version 31, <https://www.ebi.ac.uk/chembl/>), a collection of six proteins representing multiple target families (**Table 1**, column: Class) was identified using Uniprot IDs and ChEMBL target IDs. These targets include the alpha-1a adrenergic and dopamine D2 receptors as members of the GPCR family, glycogen synthase kinase-3 beta (GSK3 β) from the kinase family, Matrix-metalloprotease 3, also known as MMP3 or stromelysin as metal-loprotease, the sodium channel Nav1.7 as ion channel and the peroxisome proliferator-activated receptor delta (PPAR δ) as member of the nuclear hormone receptor (NHR) family. See Supporting Methods for details on how these datasets were derived.

Sanofi datasets

Four internal sets with structure-activity relationship (SAR) data were used. These allow us to overcome limitations of public SAR datasets which may merge assay data from different sources and which do not provide specific information on the order in which experiments were carried out. Such information allows direct comparison of the active learning solution and current best practices. The first dataset comprises compounds for inhibiting serine protease factor Xa (FXa) [Nazaré et al. \(2004\)](#); [Nazare et al. \(2012\)](#); [Matter et al. \(2002\)](#); [Nazaré et al. \(2005\)](#); [Matter et al. \(2005\)](#). The second dataset is for inhibiting the aspartyl protease Renin [Scheiper et](#)

[al. \(2010\)](#); [Matter et al. \(2011\)](#); [Scheiper et al. \(2011\)](#)) The third dataset is for the Matrix-metalloprotease 8, MMP-8, (human neutrophil collagenase) and the final dataset provides is for agonists for the nuclear hormone receptor PPAR δ (peroxisome proliferator-activated receptor δ).

Each data point in the datasets comprises the chemical structure of the molecule, represented as a Simplified Molecular Input Line Entry System (SMILES) [Weininger \(1988\)](#), and the target value, which denotes the molecular property as a scalar. **Table 1** provides details on the target and size of the each of these datasets.

Active learning for small molecule optimization

We use active learning to optimize experiment selection. We assume a setup whereby a user can select from among a pool of unlabelled samples to query for their labels, i.e., to test experimentally. Furthermore, batches of experiments are performed in rounds.

In a given round, we denote the set of possible samples to select for labelling as $\mathcal{V}$, and the batch size as B . The *ideal* active learner will select a subset, $\mathcal{B} \subset \mathcal{V}$, $|\mathcal{B}| = B$, of these experiments such that, after the experiments are performed and the labelled samples are added to the training data, upon averaging it is expected that the overall loss of the *resulting* models is minimized.

In the original active learning frameworks, only one sample is queried at a time [Lewis and Gale \(1994\)](#) [Dagan and Engelson \(1995\)](#). In this case, a straightforward and effective strategy is to query the sample on which the current model's prediction has the highest epistemic uncertainty. While this works well for single queries, it is hard to use the same approach for batch active learning. In many cases, the most uncertain queries will be similar to each other, and so their uncertainties are not independent, and so choosing as a batch the B most uncertain samples will be redundant, and not the most effective use of the queries [Azimi et al. \(2012\)](#).

In this paper we tested several different methods for selecting batches, including two new methods we developed. These two new methods are similar to “select the least certain samples” in spirit, selecting batches which maximize the total covariance of their uncertainty (estimated in different ways).

Problem Formulation: Batch Active Learning for Deep Learning Regression Models

We consider a batch active learning scenario with multiple rounds of selection, $\{1, 2, \dots, T\}$. At the t -th round, let $\mathcal{L}^{(t)}$ be the labeled dataset and $\mathcal{V}^{(t)}$ be the unlabeled dataset. To be clear, $\mathcal{V}^{(t)}$ and $\mathcal{L}^{(t)}$ are determined only after the selection method has had a chance to examine $\mathcal{V}^{(t-1)}$ and $\mathcal{L}^{(t-1)}$.

Since in our analysis we usually consider the selection problem in a *single round*, we omit superscripts $^{(t)}$ where it won't cause confusion.

Let $f_{\theta} : \mathcal{X} \rightarrow \mathbb{R}$ be the trained model, determined by its parameter $\theta \in \Theta$. As usual for regression problems, θ is chosen to approximately solve the following optimization problem,

$$\theta = \min_{\theta \in \Theta} \sum_{(x,y) \in \mathcal{L}} \ell(x, y; \theta), \quad \ell(x, y; \theta) = \frac{1}{2} \|y - f_{\theta}(x)\|_2^2 + \frac{1}{2} \lambda \|\theta\|_2^2, \quad (1)$$

where $R(\theta) = \frac{1}{2} \lambda \|\theta\|_2^2$ is the L^2 -regularization term for the weights θ . The optimization problem is (approximately) solved with a variation of SGD (in our case the popular ADAM [Kingma and Ba \(2014\)](#)), by iterating over multiple mini batches from the training set $\mathcal{L}$.

In Active Learning, a *selection method* S is a function

$$S : (\mathcal{L}, \mathcal{U}, \theta) \mapsto \mathcal{B} \subset \mathcal{U}, \quad |\mathcal{B}| = B \quad (2)$$

That is, B samples are selected from the unlabeled dataset $\mathcal{U}$ by the active learning algorithm, and the selection may use the information of the given sets $\mathcal{U}$ and $\mathcal{L}$, as well as the current supervised model f_θ , which is trained using the labeled dataset $\mathcal{L}$.

Typically, selection is done by (approximately) solving an optimization problem,

$$\max_{\substack{\mathcal{B} \subseteq \mathcal{U} \\ |\mathcal{B}| = B}} a(\mathcal{B}, \mathcal{L}, \mathcal{U}, \theta) \in \mathbb{R}$$

where the *acquisition function*, a , depends on the selection method. In a single selection round, $\mathcal{L}$, $\mathcal{U}$ and θ are given to the selector, and from the selector's point of view are fixed; therefore, in this case a depends only on the choice $\mathcal{B}$.

Predictive uncertainty

Estimating predictive uncertainty is essential to our batch selection methods. Predictive uncertainty can be broadly divided into: *epistemic* uncertainty and *aleatoric* uncertainty, which may be roughly understood as *model uncertainty* and *noise*, respectively.

Aleatoric uncertainty is predictive uncertainty due to the inability of the model to give a precise prediction, even when the model parameters are uniquely specified. If θ are the model parameters, then we denote this distribution $p_{AI}(y | x, \theta)$.

Epistemic uncertainty is predictive uncertainty due to uncertainty in the model's parameters. This is precisely the uncertainty which could be reduced by more labelled data, and therefore is the relevant quantity for active learning. In the strict Bayesian setting, if $f_\theta : \mathcal{X} \rightarrow \mathbb{R}$ is the model corresponding to parameters θ , and $p(\theta)$ is the posterior distribution given the training data $\mathcal{D}$, this distribution is

$$p(y | x, D) = \int p(y | x, \theta) p(\theta | D) d\theta \quad (3)$$

$$= \mathbb{E}_{p(\theta | D)} [p(y | x, \theta)] \quad (4)$$

Calculating this distribution exactly is intractable, and so we use approximations, described in Section.

Batch selection via determinant maximization

In the batch setting, selecting the batch with the top-ranked epistemic uncertainty may lead to a highly correlated selection set, and thus wasted experiments. I.e., if a sample x has the highest estimated variance, then probably very small perturbations of x will have similarly high variance, and a variance-only selection method will give a high score to this batch of similar (and thus correlated) samples.

We thus seek to select the batch with the highest total independent uncertainty. This can be computed by selecting the batch with the highest joint entropy, which, under the assumption of a normal distribution, is the highest log-determinant of the (epistemic) covariance.

$$\operatorname{argmax}_{\mathcal{B} \subset \mathcal{U}, |\mathcal{B}| = N} \log(\det(\operatorname{cov}(\mathcal{B}))) \quad (5)$$

Even if we already know the epistemic covariance for any batch (a problem we address in the next section), finding the strict maximum is NP-hard [Ohsaka \(2022\)](#). Therefore we use an approximate maximization technique: We randomly generate a collection of batches as starting points, each containing N distinct samples independently chosen from a distribution proportional to the quantile of the variances. Then we select the best $M < N$ of these batches as starting points for optimization. Then, for each starting point, we optimize the batch element-wise, i.e., changing the first element to optimize the covariance, then changing the second, and so on, doing several passes until the process reaches equilibrium. Then, we select the highest-scoring final batch.

Approximation of the posterior distribution

A straightforward way to get a distribution of predictions is to train an ensemble of models [Lakshminarayanan et al. \(2017\)](#), and to use the outputs of these models to give an ensemble of predictions. However, this approach involves multiple rounds of retraining, which is resource intensive. Furthermore, there is no guarantee that the ensemble of trained models will be diverse, and no guarantee that it will approximate the true Bayesian posterior.

Instead, we take a more economical approach by leveraging only one trained model. The idea behind this method is that the optimal parameters of a trained deep regression model, i.e. θ^* , are the *maximum a posteriori* (MAP) estimation of an equivalent Bayesian deep regression problem.

Thus, the Bayesian inference of the posterior of θ is approximated by leveraging the computed MAP estimation of θ^* . As it is shown in [Eq. \(1\)](#), the optimal θ^* could be treated as the MAP estimation of the probabilistic model of $\{Y, \theta\}$ where $Y = \{y\}_{y \in \mathcal{L}}$.

$$\begin{aligned} p(Y, \theta) &= p(\theta)p(Y|\theta) \\ p(\theta) &\propto \exp\left(-\frac{1}{2}R(\theta)\right) \\ p(Y|\theta) &= \prod_{x, y \in \mathcal{L}} \mathcal{N}(y|f_{\theta}(x), 1) \end{aligned} \quad (6)$$

Two approximations for computing the epistemic covariance

We developed and tested two different methods to approximate the covariance:

- Monte Carlo Dropout

The first approach we developed is based on Monte Carlo dropout. Dropout [Hinton et al. \(2012\)](#) is a well-known technique for better training of neural networks. With dropout, in the back propagation stage of the optimization, each of the neural network's weights are randomly set to 0 with a certain probability (as defined by the drop out ratio, r). Previous work has demonstrated that models built with dropout are less prone to overfitting and that training with dropout is quite similar to variational inference of the probabilistic models [Gal and Ghahramani \(2016\)](#); [Lawrence \(2001\)](#) when assuming the following form of the posterior approximation $q(\theta)$,

$$\theta = \theta^* \cdot M \quad (7)$$

$$M \sim \text{Bernoulli}(r) \quad (8)$$

where r is the dropout ratio, M , the masks, are the $\{0, 1\}$ variables of same shape as θ , and θ^* are the weights obtained after training with dropout. We thus use dropout to obtain S predictions by sampling S masks $\{M_s\}_{s \in 1, 2, \dots, S}$ according to [Eq. \(8\)](#). We then use the new network to predict sample labels and construct the covariance matrix based on these predictions.

- Laplace Approximation

In addition to dropout, we also employed the *Laplace approximation* for estimating the posterior. This method assumes that the posterior of the model parameters is a multi-variate normal distribution centered at the MAP value, and with variance equal to the Fisher information, which can be found through a straightforward calculation. This simplifying assumption makes the approximation fast to use, as the computation relies only on MAP estimation and differentiation of the loss function.

We assume the following for the posterior probability distribution, $q(\theta)$:

$$q(\theta) = \mathcal{N}(\theta | \theta^*, \Lambda^{*-1}) \quad (9)$$

By matching 0th and 2nd derivatives, we find Λ , and thus our covariance matrix Λ^{*-1} . The covariance matrix Λ^{*-1} does not have to be represented as a full matrix of shape $Q \times Q$, which is usually very large for neural networks. For example, Laplace Redux [Daxberger et al. \(2021\)](#) further approximate the multi-variate normal distribution in Eq. (9) significantly reducing the computation cost. See Supporting Methods for additional details.

Comparisons to other methods

To compare our proposed approach with prior approaches, for batch active learning we looked at three previous methods suggested for this task: Random selection [Settles \(2009\)](#), optimizing uncertainty and diversity based on information theory (the *BAIT* method), and an unsupervised method that is optimizing only for diversity [Nguyen and Smeulders \(2004\)](#) (a method we term *k-means sampling*). See Supporting Methods for details on these prior methods and how we used them here.

Evaluation Experiments

For the retrospective experiment an existing labelled data set is selected to simulate an active learning experiment. We start by selecting a random subset (which is used as the initial set for all comparisons as well).

For our data the chemical structures in X are represented as molecular graphs using the MolGraphConvFe from the DeepChem library [Kearnes et al. \(2016\)](#). For each active learning cycle, models are trained using DeepChem [dee \(2016\)](#), a library which provides the implementation of the Graph Convolutional Neural Networks (GCNN) for molecular systems [Kipf and Welling \(2016\)](#).

For accuracy as well as model performance we use the root mean squared error (RMSE). See Supporting Methods for details on how the RMSE is computed.

Neural network architecture

For all our datasets we use the “neural fingerprints” class of models described in [Duvenaud et al. \(2015\)](#), as implemented in the DeepChem library GraphConvModel [contributors \(2022\)](#). See Supporting Methods for complete details. Although our method is compatible with, PyTorch, TensorFlow, and Keras frameworks, the benchmarks presented here are performed using the Keras framework within Deepchem suite. In order to enforce deterministic behaviour of the models for each selection methods and active learning rounds, during the training the model weights were manually loaded at the beginning of each active learning loop. We perform multiple runs of the retrospective experiments for each methods.

Code and Data availability

Code and all data used in this study are available as an open source package at Sanofi-Public GitHub page.

References

(2016) **Democratizing Deep-Learning for Drug Discovery, Quantum Chemistry, Materials Science and Biology**. GitHub; 2016. <https://github.com/deepchem/deepchem>.

Ash J, Goel S, Krishnamurthy A, Kakade S., Ranzato M, Beygelzimer A, Dauphin Y, Liang PS, Vaughan JW (2021) **Gone Fishing: Neural Active Learning with Fisher Embeddings** *Advances in Neural Information Processing Systems* :8927–8939

Azimi J, Fern A, Zhang-Fern X, Borradaile G, Heeringa B. (2012) **Batch active learning via coordinated matching** *arXiv*

Cai C, Wang S, Xu Y, Zhang W, Tang K, Ouyang Q, Lai L, Pei J. (2020) **Transfer Learning for Drug Discovery** *Journal of Medicinal Chemistry* **63**:8683–8694 <https://doi.org/10.1021/acs.jmedchem.9b02147>

Chuang KV, Gunsalus LM, Keiser MJ (2020) **Learning Molecular Representations for Medicinal Chemistry** *Journal of Medicinal Chemistry* **63**:8705–8722 <https://doi.org/10.1021/acs.jmedchem.0c00385>

Cohn D, Atlas L, Ladner R. (1994) **Improving generalization with active learning** *Machine learning* **15**:201–221

Cohn DA, Ghahramani Z, Jordan MI (1996) **Active learning with statistical models** *Journal of artificial intelligence research* **4**:129–145

Cohn DA, Ghahramani Z, Jordan MI (1996) **Active learning with statistical models** *Journal of artificial intelligence research* **4**:129–145

(2022) **contributors D, DeepChem Documentation - Keras Models - GraphConvModel; 2022.** https://deepchem.readthedocs.io/en/latest/api_reference/models.html#graphconvmodel, x[Online; accessed 27-February-2023].

Dagan I, Engelson SP, Prieditis A, Russell S (1995) **Committee-Based Sampling For Training Probabilistic Classifiers** *Machine Learning Proceedings 1995* :150–157 <https://doi.org/10.1016/B978-1-55860-377-6.50027-X>

Daxberger E, Kristiadi A, Immer A, Eschenhagen R, Bauer M, Hennig P. (2021) **Laplace Redux—Effortless Bayesian Deep Learning** *In: NeurIPS*

Ding X, Cui R, Yu J, Liu T, Zhu T, Wang D, Chang J, Fan Z, Liu X, Chen K, Jiang H, Li X, Luo X, Zheng M. (2021) **Active Learning for Drug Design: A Case Study on the Plasma Exposure of Orally Administered Drugs** *Journal of Medicinal Chemistry* **64**:16838–16853 <https://doi.org/10.1021/acs.jmedchem.1c01683>

Duvenaud DK, Maclaurin D, Iparraguirre J, Bombarell R, Hirzel T, Aspuru-Guzik A, Adams RP, Cortes C, Lawrence N, Lee D, Sugiyama M, Garnett R (2015) **Convolutional Networks on Graphs for Learning Molecular Fingerprints** *Advances in Neural Information Processing Systems*

- Gal Y, Ghahramani Z., of Proceedings of Machine Learning Research, Balcan MF, Weinberger KQ (2016) **Dropout as a Bayesian Approximation: Representing Model Uncertainty in Deep Learning** *Proceedings of The 33rd International Conference on Machine Learning* :1050–1059
- Gaulton A, Bellis LJ, Bento AP, Chambers J, Davies M, Hersey A, Light Y, McGlinchey S, Michalovich D, Al-Lazikani B, et al. (2012) **ChEMBL: a large-scale bioactivity database for drug discovery** *Nucleic acids research* **40**:D1100–D1107
- Grebner C, Matter H, Hessler G., Heifetz A (2022) **Artificial Intelligence in Compound Design** :349–382 https://doi.org/10.1007/978-1-0716-1787-8_15
- Haghighatlari M, Vishwakarma G, Altarawy D, Subramanian R, Kota BU, Sonpal A, Setlur S, Hachmann J. (2019) **ChemML: A Machine Learning and Informatics Program Package for the Analysis, Mining, and Modeling of Chemical and Materials Data** *ChemRxiv* <https://doi.org/10.26434/chemrxiv.8323271.v1>
- Hessler G, Baringhaus KH (2018) **Artificial Intelligence in Drug Design** *Molecules* **23** <https://doi.org/10.3390/molecules23102520>
- Hinton GE, Srivastava N, Krizhevsky A, Sutskever I, Salakhutdinov RR (2012) **Improving neural networks by preventing co-adaptation of feature detectors** *arXiv*
- Huang K, Fu T, Gao W, Zhao Y, Roohani Y, Leskovec J, Coley CW, Xiao C, Sun J, Zitnik M. (2022) **Artificial intelligence foundation for therapeutic science** *Nature Chemical Biology* **18**:1033–1036 <https://doi.org/10.1038/s41589-022-01131-2>
- Kearnes S, McCloskey K, Berndl M, Pande V, Riley P. (2016) **Molecular graph convolutions: moving beyond finger-prints** *Journal of Computer-Aided Molecular Design* **30**:595–608 <https://doi.org/10.1007/s10822-016-9938-8>
- Kendall A, Gal Y. (2017) **What uncertainties do we need in bayesian deep learning for computer vision?** *Advances in neural information processing systems* **30**
- Kingma DP, Ba J. (2014) **Adam: A method for stochastic optimization** *arXiv*
- Kipf TN, Welling M. (2016) **Semi-Supervised Classification with Graph Convolutional Networks**
- Lakshminarayanan B, Pritzel A, Blundell C. (2017) **Simple and scalable predictive uncertainty estimation using deep ensembles** *Advances in neural information processing systems* **30**
- Lawrence ND (2001) **Lawrence ND. Variational Inference in Probabilistic Models.** University of Cambridge; 2001. <https://books.google.de/books?id=mbi5rQEACAAJ>.
- Lewis DD, Gale WA, Croft BW, van Rijsbergen CJ (1994) **A Sequential Algorithm for Training Text Classifiers** *SIGIR '94* :3–12
- Matter H, Defossa E, Heinelt U, Blohm PM, Schneider D, Müller A, Herok S, Schreuder H, Liesum A, Brachvogel V, Lönze P, Walser A, Al-Obeidi F, Wildgoose P. (2002) **Design and Quantitative Structure-Activity Relationship of 3-Amidinobenzyl-1H-indole-2-carboxamides as Potent, Nonchiral, and Selective Inhibitors of Blood Coagulation Factor Xa** *J Med Chem* **45**:2749–2769

- Matter H, Scheiper B, Steinhagen H, Böcskei Z, Fleury V, McCort G. (2011) **Structure-based design and optimization of potent renin inhibitors on 5- or 7-azaindole-scaffolds** *Bioorganic & Medicinal Chemistry Letters* **21**:5487–5492 <https://doi.org/10.1016/j.bmcl.2011.06.112>
- Matter H, Will DW, Nazaré M, Schreuder H, Laux V, Wehner V. (2005) **Structural Requirements for Factor Xa Inhibition by 3-Oxybenzamides with Neutral P1 Substituents: Combining X-ray Crystallography, 3D-QSAR, and Tailored Scoring Functions** *Journal of Medicinal Chemistry* **48**:3290–3312 <https://doi.org/10.1021/jm0491871>
- Mehrjou A, Soleymani A, Jesson A, Notin P, Gal Y, Bauer S, Schwab P. (2021) **GeneDisco: A Benchmark for Experimental Design in Drug Discovery**
- Nazaré M, Essrich M, Will DW, Matter H, Ritter K, Urmann M, Bauer A, Schreuder H, Dudda A, Czech J, et al. (2004) **Factor Xa inhibitors based on a 2-carboxyindole scaffold: SAR of neutral P1 substituents** *Bioorganic & medicinal chemistry letters* **14**:4191–4195
- Nazare M, Matter H, Will DW, Wagner M, Urmann M, Czech J, Schreuder H, Bauer A, Ritter K, Wehner V. (2012) **Fragment Deconstruction of Small, Potent Factor Xa Inhibitors: Exploring the Superadditivity Energetics of Fragment Linking in Protein-Ligand Complexes** *Angewandte Chemie International Edition* **51**:905–911 <https://doi.org/10.1002/anie.201107091>
- Nazaré M, Will DW, Matter H, Schreuder H, Ritter K, Urmann M, Essrich M, Bauer A, Wagner M, Czech J, Lorenz M, Laux V, Wehner V. (2005) **Probing the Subpockets of Factor Xa Reveals Two Binding Modes for Inhibitors Based on a 2-Carboxyindole Scaffold: A Study Combining Structure-Activity Relationship and X-ray Crystallography** *Journal of Medicinal Chemistry* **48**:4511–4525 <https://doi.org/10.1021/jm0490540>
- Nguyen HT, Smeulders A. (2004) **Active Learning Using Pre-Clustering** *In: NeurIPS ICML 's04* <https://doi.org/10.1145/1015330.1015349>
- Ohsaka N (2022) **On the Parameterized Intractability of Determinant Maximization** <https://doi.org/10.48550/ARXIV.2209.12519>
- Pertusi DA, Moura ME, Jeffries JG, Prabhu S, Walters Biggs B, Tyo KEJ (2017) **Predicting novel substrates for enzymes with minimal experimental effort with active learning** *Metabolic Engineering* **44**:171–181 <https://doi.org/10.1016/j.ymben.2017.09.016>
- Reker D, Schneider P, Schneider G, Brown J. (2017) **Active learning for computational chemogenomics** *Future Medicinal Chemistry* **9**:381–402 <https://doi.org/10.4155/fmc-2016-0197>
- Scheiper B, Matter H, Steinhagen H, Böcskei Z, Fleury V, McCort G. (2011) **Structure-based optimization of potent 4- and 6-azaindole-3-carboxamides as renin inhibitors** *Bioorganic & Medicinal Chemistry Letters* **21**:5480–5486 <https://doi.org/10.1016/j.bmcl.2011.06.114>
- Scheiper B, Matter H, Steinhagen H, Stilz U, Böcskei Z, Fleury V, McCort G. (2010) **Discovery and optimization of a new class of potent and non-chiral indole-3-carboxamide-based renin inhibitors** *Bioorganic & Medicinal Chemistry Letters* **20**:6268–6272 <https://doi.org/10.1016/j.bmcl.2010.08.092>
- Settles B. (2009) **Active Learning Literature Survey**

- Settles B. (2012) **Active Learning** *Synthesis Lectures on Artificial Intelligence and Machine Learning*
- Sorkun MC, Khetan A, Er S. (2019) **AqSolDB, a curated reference set of aqueous solubility and 2D descriptors for a diverse set of compounds** *Scientific Data* **6** <https://doi.org/10.1038/s41597-019-0151-1>
- Thompson J, Walters WP, Feng JA, Pabon NA, Xu H, Goldman BB, Moustakas D, Schmidt M, York F. (2022) **Optimizing active learning for free energy calculations** *Artificial Intelligence in the Life Sciences* **2** <https://doi.org/10.1016/j.ailsci.2022.100050>
- Tynes M, Gao W, Burrill DJ, Batista ER, Perez D, Yang P, Lubbers N. (2021) **Pairwise Difference Regression: A Machine Learning Meta-algorithm for Improved Prediction and Uncertainty Quantification in Chemical Search** *Journal of Chemical Information and Modeling* **61**:3846–3857 <https://doi.org/10.1021/acs.jcim.1c00670>
- Wang NN, Dong J, Deng YH, Zhu MF, Wen M, Yao ZJ, Lu AP, Wang JB, Cao DS (2016) **ADME Properties Evaluation in Drug Discovery: Prediction of Caco-2 Cell Permeability Using a Combination of NSGA-II and Boosting** *Journal of Chemical Information and Modeling* **56**:763–773 <https://doi.org/10.1021/acs.jcim.5b00642>
- Weininger D. (1988) **SMILES, a chemical language and information system. 1. Introduction to methodology and encoding rules** *Journal of Chemical Information and Computer Sciences* **28**:31–36 <https://doi.org/10.1021/ci00057a005>
- Weiss K, Khoshgoftaar TM, Wang D. (2016) **A survey of transfer learning** *Journal of Big Data* **3** <https://doi.org/10.1186/s40537-016-0043-6>
- Wenlock M, Tomkinson N (2015) **Experimental in vitro DMPK and physicochemical data on a set of publicly disclosed compounds**
- Wu F, Zhou Y, Li L, Shen X, Chen G, Wang X, Liang X, Tan M, Huang Z. (2020) **Computational Approaches in Preclinical Studies on Drug Discovery and Development** *Frontiers in Chemistry* **8** <https://doi.org/10.3389/fchem.2020.00726>
- Wu Z, Ramsundar B, Feinberg EN, Gomes J, Geniesse C, Pappu AS, Leswing K, Pande V. (2018) **MoleculeNet: a benchmark for molecular machine learning** *Chemical science* **9**:513–530
- Xiong G, Wu Z, Yi J, Fu L, Yang Z, Hsieh C, Yin M, Zeng X, Wu C, Lu A, Chen X, Hou T, Cao D. (2021) **ADMETlab 2.0: an integrated online platform for accurate and comprehensive predictions of ADMET properties** *Nucleic Acids Research* **49**:W5–W14 <https://doi.org/10.1093/nar/gkab255>
- Zaverkin V, Holzmüller D, Steinwart I, Kästner J. (2022) **Exploring chemical and conformational spaces by batch mode deep active learning** *Digital Discovery* **1**:605–620 <https://doi.org/10.1039/D2DD00034B>

Author information

Michael Bailey

R&D Data & Computational Science, Sanofi, Cambridge, MA, United States

For correspondence: sven.jager@sanofi.com

ORCID iD: [0009-0004-8425-993X](https://orcid.org/0009-0004-8425-993X)

Saeed Moayedpour

R&D Data & Computational Science, Sanofi, Cambridge, MA, United States

Ruijiang Li

Digital Data, Sanofi, Shanghai, China

Alejandro Corrochano-Navarro

R&D Data & Computational Science, Sanofi, Cambridge, MA, United States

Alexander Kötter

Synthetic Molecular Design, Integrated Drug Discovery, Sanofi-Aventis Deutschland GmbH, Industriepark Höchst, Building G838, 65926 Frankfurt am Main, Germany

Lorenzo Kogler-Anele

R&D Data & Computational Science, Sanofi, Cambridge, MA, United States

Saleh Riahi

R&D Data & Computational Science, Sanofi, Cambridge, MA, United States

Christoph Grebner

Synthetic Molecular Design, Integrated Drug Discovery, Sanofi-Aventis Deutschland GmbH, Industriepark Höchst, Building G838, 65926 Frankfurt am Main, Germany

Gerhard Hessler

Synthetic Molecular Design, Integrated Drug Discovery, Sanofi-Aventis Deutschland GmbH, Industriepark Höchst, Building G838, 65926 Frankfurt am Main, Germany

Hans Matter

Synthetic Molecular Design, Integrated Drug Discovery, Sanofi-Aventis Deutschland GmbH, Industriepark Höchst, Building G838, 65926 Frankfurt am Main, Germany

Marc Bianciotto

Molecular Design Sciences, Integrated Drug Discovery, Sanofi, Vitry-sur-Seine, 94403, France

Pablo Mas

Molecular Design Sciences, Integrated Drug Discovery, Sanofi, Vitry-sur-Seine, 94403, France

Ziv Bar-Joseph

R&D Data & Computational Science, Sanofi, Cambridge, MA, United States

For correspondence: ziv.bar-joseph@sanofi.com

ORCID iD: [0000-0003-3430-6051](https://orcid.org/0000-0003-3430-6051)

Sven Jager

R&D Data & Computational Science, Sanofi, Cambridge, MA, United States

For correspondence: sven.jager@sanofi.com

ORCID iD: [0000-0003-1410-9114](https://orcid.org/0000-0003-1410-9114)

Editors

Reviewing Editor

Alan Talevi

National University of La Plata, Argentina

Senior Editor

Volker Dötsch

Goethe University, Germany

Reviewer #1 (Public Review):

The authors present a study focused on addressing the key challenge in drug discovery, which is the optimization of absorption and affinity properties of small molecules through in silico methods. They propose active learning as a strategy for optimizing these properties and describe the development of two novel active learning batch selection methods. The methods are tested on various public datasets with different optimization goals and sizes, and new affinity datasets are curated to provide up-to-date experimental information. The authors claim that their active learning methods outperform existing batch selection methods, potentially reducing the number of experiments required to achieve the same model performance. They also emphasize the general applicability of their methods, including compatibility with popular packages like DeepChem.

Strengths:

Relevance and Importance: The study addresses a significant challenge in the field of drug discovery, highlighting the importance of optimizing the absorption and affinity properties of small molecules through in silico methods. This topic is of great interest to researchers and pharmaceutical industries.

Novelty: The development of two novel active learning batch selection methods is a commendable contribution. The study also adds value by curating new affinity datasets that provide chronological information on state-of-the-art experimental strategies.

Comprehensive Evaluation: Testing the proposed methods on multiple public datasets with varying optimization goals and sizes enhances the credibility and generalizability of the findings. The focus on comparing the performance of the new methods against existing batch selection methods further strengthens the evaluation.

Weaknesses:

Lack of Technical Details: The feedback lacks specific technical details regarding the developed active learning batch selection methods. Information such as the underlying algorithms, implementation specifics, and key design choices should be provided to enable readers to understand and evaluate the methods thoroughly.

Evaluation Metrics: The feedback does not mention the specific evaluation metrics used to assess the performance of the proposed methods. The authors should clarify the criteria employed to compare their methods against existing batch selection methods and demonstrate the statistical significance of the observed improvements.

Reproducibility: While the authors claim that their methods can be used with any package, including DeepChem, no mention is made of providing the necessary code or resources to reproduce the experiments. Including code repositories or detailed instructions would enhance the reproducibility and practical utility of the study.

Suggestions for Improvement:

Elaborate on the Methodology: Provide an in-depth explanation of the two active learning batch selection methods, including algorithmic details, implementation considerations, and any specific assumptions made. This will enable readers to better comprehend and evaluate the proposed techniques.

Clarify Evaluation Metrics: Clearly specify the evaluation metrics employed in the study to measure the performance of the active learning methods. Additionally, conduct statistical tests to establish the significance of the improvements observed over existing batch selection methods.

Enhance Reproducibility: To facilitate the reproducibility of the study, consider sharing the code, data, and resources necessary for readers to replicate the experiments. This will allow researchers in the field to validate and build upon your work more effectively.

Conclusion:

The authors' study on active learning methods for optimizing drug discovery presents an important and relevant contribution to the field. The proposed batch selection methods and curated affinity datasets hold promise for improving the efficiency of drug discovery processes. However, to strengthen the study, it is crucial to provide more technical details, clarify evaluation metrics, and enhance reproducibility by sharing code and resources. Addressing these limitations will further enhance the value and impact of the research.

Reviewer #2 (Public Review):

The authors presented a well-written manuscript describing the comparison of active-learning methods with state-of-art methods for several datasets of pharmaceutical interest. This is a very important topic since active learning is similar to a cyclic drug design campaign such as testing compounds followed by designing new ones which could be used to further tests and a new design cycle and so on. The experimental design is comprehensive and adequate for proposed comparisons. However, I would expect to see a comparison regarding other regression metrics and considering the applicability domain of models which are two essential topics for the drug design modelers community.

Author Response:

We thank the reviewers for their constructive comments. Below we include a point by point response.

Reviewer #1 (Public Review):

[...] Elaborate on the Methodology: Provide an in-depth explanation of the two active learning batch selection methods, including algorithmic details, implementation considerations, and any specific assumptions made. This will enable readers to better comprehend and evaluate the proposed techniques.

We thank the reviewer for this suggestion. Following this comments we will extend the text in Methods (in Section: Batch selection via determinant maximization and Section: Approximation of the posterior distribution) and in Supporting Methods (Section: Toy example). We will also include the pseudo code for the Batch optimization method.

Clarify Evaluation Metrics: Clearly specify the evaluation metrics employed in the study to measure the performance of the active learning methods. Additionally, conduct statistical tests to establish the significance of the improvements observed over existing batch selection methods.

Following this comment we will add to Table 1 details about the way we computed the cutoff times for the different methods. We will also provide more details on the statistics we performed to determine the significance of these differences.

Enhance Reproducibility: To facilitate the reproducibility of the study, consider sharing the code, data, and resources necessary for readers to replicate the experiments. This will allow researchers in the field to validate and build upon your work more effectively.

This is something we already included with the original submission. The code is publicly available. In fact, we provide a python library, ALIEN (Active Learning in data Exploration) which is published on the Sanofi Github (<https://github.com/Sanofi-Public/Alien>). We also provide details on the public data used and expect to provide the internal data as well. We included a small paragraph on code and data availability.

Reviewer #2 (Public Review):

[...] I would expect to see a comparison regarding other regression metrics and considering the applicability domain of models which are two essential topics for the drug design modelers community.

We want to thank the reviewer for these comments. We will provide a detailed response to their specific comments when we resubmit.