



Universiteit
Leiden
The Netherlands

To explore drug space smarter: artificial intelligence in drug design for G protein-coupled receptors

Liu, X.

Citation

Liu, X. (2022, February 15). *To explore drug space smarter: artificial intelligence in drug design for G protein-coupled receptors*. Retrieved from <https://hdl.handle.net/1887/3274010>

Version: Publisher's Version

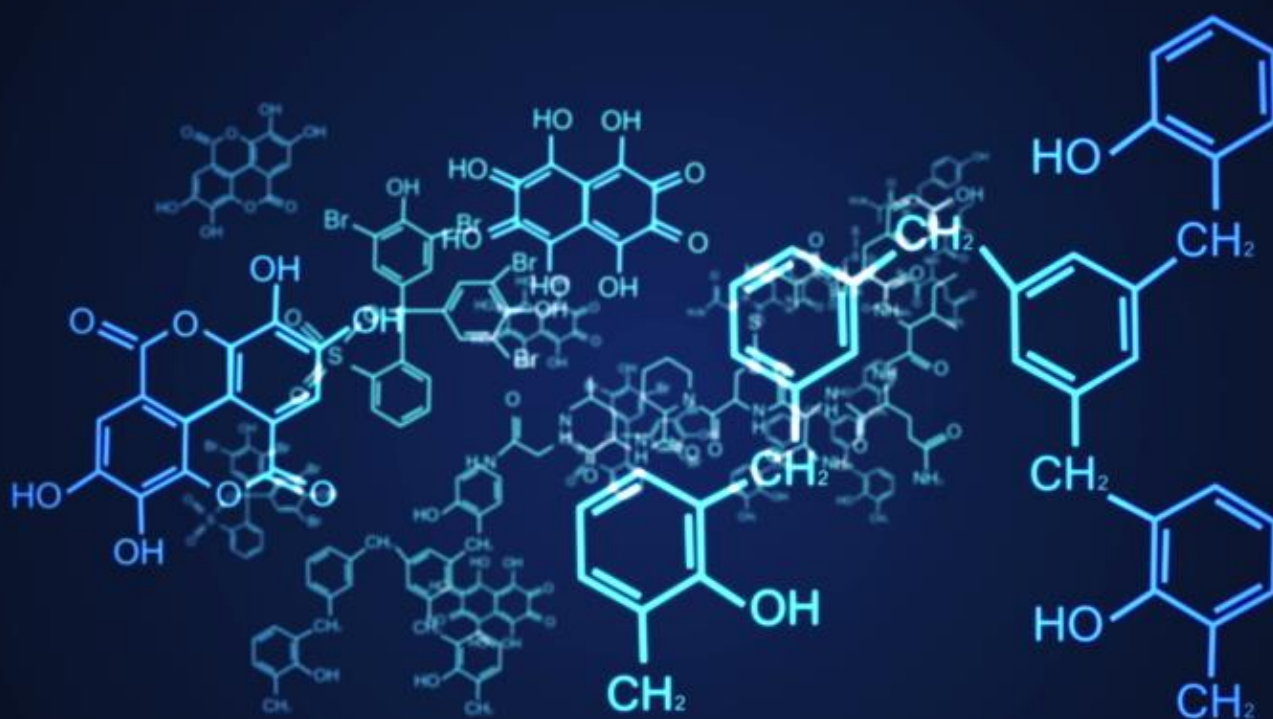
License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/3274010>

Note: To cite this publication please use the final published version (if applicable).

Chapter 5

DrugEx v3: scaffold-constrained drug design with graph Transformer-based reinforcement learning



Xuhan Liu, Kai Ye, Herman W. T. van Vlijmen, Adriaan P. IJzerman and Gerard J. P. van Westen*. Preprint. <https://doi.org/10.26434/chemrxiv-2021-px6kz>

Abstract

Due to the large drug-like chemical space available to search for feasible drug-like molecules, rational drug design often starts from the specific scaffold to which side chains/substituents are added or modified. With the rapid growth of the application of deep learning in drug discovery, a variety of effective approaches have been developed for de novo drug design. In previous work, we proposed a method named *DrugEx*, which can be applied in polypharmacology based on multi-objective deep reinforcement learning. However, the previous version is trained under fixed objectives similar to other known methods and does not allow users to input any prior information. In order to improve the general applicability, we updated *DrugEx* to design drug molecules based on the scaffold which can contain multiple fragments provided by users. In this work, the Transformer model was employed to generate the structure of molecules. The Transformer is a multi-head self-attention deep learning model containing an encoder for receiving scaffolds as input and a decoder generating molecules as output. In order to deal with the graph representation of molecules, we proposed a novel positional encoding for each atom and bond based on an adjacency matrix to extend the architecture of the Transformer. Each molecule was generated by growing and connecting procedures for the fragments in the given scaffold that were unified into one model. Moreover, we trained this generator under a reinforcement learning framework to increase the number of desired ligands. As a proof of concept, our proposed method was applied to design ligands for the adenosine A_{2A} receptor (A_{2A}AR) and compared it with SMILES-based methods. The results demonstrated its effectiveness in that 100% of generated molecules are valid and most of them had high predicted affinity value towards A_{2A}AR with given scaffold.

Keywords: deep learning, reinforcement learning, policy gradient, drug design, Transformer, multi-objective optimization

5.1. Introduction

Due to the large drug-like chemical space (*i.e.* estimated at 10^{33} - 10^{60} organic molecules) [1], it is impossible to screen every corner of it to discover optimal drug candidates, although high-throughput screening (HTS) technology has been improved significantly in recent years [2]. Commonly, the specific scaffolds derived from endogenous substances are taken as a starting point to design analogs after side chains/substituents are added or modified [3]. These fragments are used as ‘building blocks’ to develop proper drug leads with combinatorial chemistry such as growing, linking and merging [4]. After a promising drug lead has been discovered, it is further optimized by modifying side chains to improve potency and selectivity which in turn can improve safety and tolerability [5].

The adenosine receptors (ARs) belong to a class of rhodopsin-like GPCRs including four subtypes (A_1 , A_{2A} , A_{2B} and A_3). Each of them has a unique pharmacological profile, tissue distribution, and effector coupling [6,7]. ARs are ubiquitously distributed throughout the human tissues, and involved in many biological processes and diseases [8]. Because adenosine is the endogenous agonist of ARs, a number of known ligands of the ARs are adenosine analogs and have a common scaffold. Examples include purines, xanthines, triazines, pyrimidines, and the inclusion of a ribose moiety [9]. In scaffold-based rational drug design, it is generally accepted that a chemical space consisting of 10^9 diverse molecules can be sampled with only 10^3 fragments [10].

Based on rapid developments in the last decade, deep learning has achieved a breakthrough in visual recognition, natural language processing, and other data-rich fields [11]. In drug discovery, deep learning methods have also been extensively used for drug *de novo* design [12]. For distribution-directed issues, Gomez-Bombarelli *et al.* implemented variational autoencoders (VAE) to map molecules into a latent space where each point can also be decoded into unique molecules inversely [13]. They used recurrent neural networks (RNNs) to successfully learn SMILES (simplified molecular-input line-entry system) grammar and construct a distribution of molecular libraries [14]. For goal-directed issues, Sanchez-Lengeling *et al.* combined reinforcement learning and generative adversarial networks

(GANs) to develop an approach named *ORGANIC* to design active compounds toward given targets [15]. Olivecrona *et al.* proposed the *REINVENT* algorithm which updated the reinforcement learning with a Bayesian approach and combined RNNs to generate SMILES-based desired molecules [16,17]. Moreover, Lim *et al.* proposed a method for scaffold-based molecular design with a graph generative model [18]. Li *et al.* also used deep learning to develop a tool named *DeepScaffold* for this issue [19]. Arús-Pous *et al.* employed RNNs to develop a SMILES-based scaffold decorator for *de novo* drug design [20]. Yang *et al.* used the Transformer model [21] to develop a tool named *SyntaLinker* for automatic fragment linking [22].

In previous studies, we investigated the performance of RNNs and proposed a method named *DrugEx* by integrating reinforcement learning to balance distribution-directed and goal-directed tasks [23]. Furthermore, we updated it with multi-objective reinforcement learning and applied it in polypharmacology [24]. However, the well-trained model cannot receive any input data from users and only reflect the distribution of the desired molecules with fixed conditions. If the objectives are changed, the model needs to be trained again. In this work, we compared different end-to-end deep learning methods and updated the *DrugEx* model to allow users to provide prior information, such as fragments that should occur in the generated molecules. Based on the extensive experience in our group with the A_{2A}AR, we continue to take this target as an example to evaluate the performance of our proposed methods. In the following context, we will discuss the case of scaffold-constrained drug design, *i.e.* the model takes the scaffolds containing multiple fragments as input to generate desired molecules which also can be predicted to be active to A_{2A}AR. All python code for this study is freely available at <http://github.com/XuhanLiu/DrugEx>.

5.2. Materials and methods

5.2.1. Data source

Chemical compounds were downloaded from ChEMBL using a SMILES notation (version 27) [25]. After data preprocessing implemented by RDKit, which included neutralizing charges, removing metals and small fragments, ~1.7 million molecules remained for model

pre-training. These data were reused from the work about *DrugEx v2* (*ChEMBL* set) [24]. In addition, 10,828 ligands and bioactivity data were extracted from ChEMBL to construct the *LIGAND* set containing structures and activities from bioassays towards four human adenosine receptors. The *LIGAND* set was used for fine-tuning the generative model. Molecules with annotated A_{2A}AR activity were used to train a prediction model. If multiple measurements for the same ligands existed, the average pChEMBL value (pX, including pK_i, pK_d, pIC₅₀ or pEC₅₀) was calculated and duplicate items were removed. In order to judge if the molecule is desired or not, the threshold of affinity was defined as pX = 6.5 to predict if the compound was active (≥ 6.5) or inactive (< 6.5).

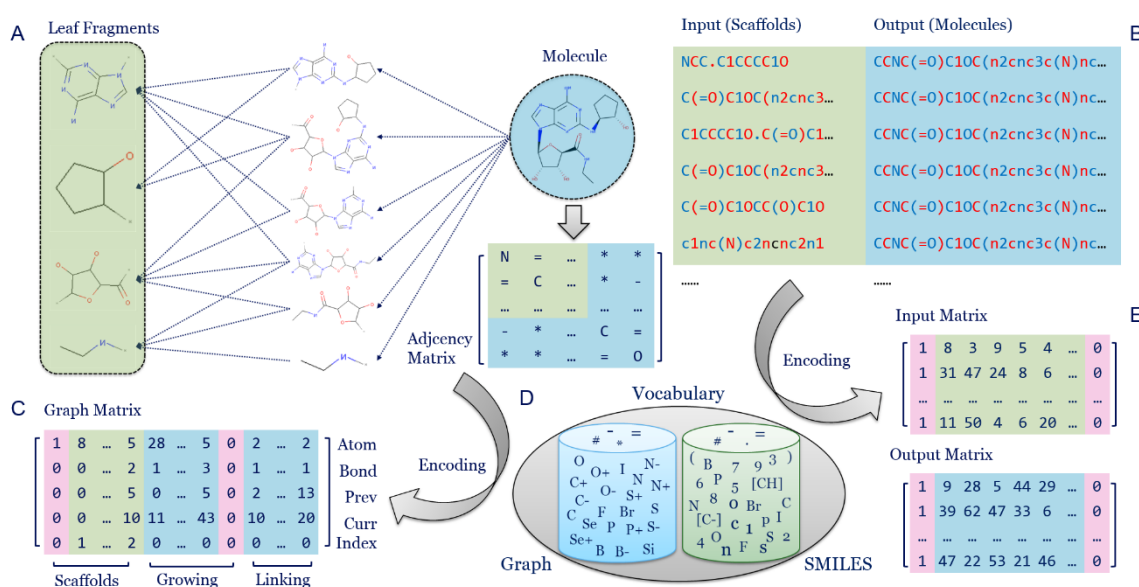


Fig. 5.1: scaffold-molecule pair dataset construction. (A) Each molecule in the dataset is decomposed hierarchically into a series of fragments with the BRICS algorithm. (B) Subsequently data pairs between input and output are created. Combinations of leaf fragments form the scaffold as input, the whole molecule becomes the output. Each token in SMILES sequences is separated by different colors. (C) After conversion to the adjacency matrix, each molecule was represented as a graph matrix. The graph matrix contains five rows, standing for the atom, bond, previous and current positions and fragment index. The columns are composed with three parts to store the information of scaffolds, growing section and linking section. (D) All of tokens are collected to construct the vocabularies for SMILES-based and graph-based generators, respectively. (E) An example of the input and output matrices for the SMILES representation of scaffolds and molecules

Furthermore, the dataset was constructed with an input-output pair for each data point. Each molecule was decomposed into a batch of fragments with BRICS methods [26] in RDKit (Fig. 5.1A). If the molecule contained more than four leaf fragments, the smaller

fragments were ignored and a maximum of four larger fragments were reserved to be randomly combined at one time. Here, the scaffold was defined as the combination of different fragments which can be either continuously (linked) or discretely (separated). Their SMILES sequences were joined with ‘.’ as input data which were paired with the full SMILES of molecules. The resulting fragments-molecule pair forms the output data (Fig. 5.1B). After completion of constructing the data pairs, the set was split into a training set and test set with the ratio 9:1 based on the input scaffolds. The resulting *ChEMBL* set contained 10,418,681 and 1,083,271 pairs for training and test set, respectively. The *LIGAND* set contained 61,413 pairs in the training set and 7,525 pairs in the test set.

5.2.2. Molecular representations

In this study we tested two different molecular representations: SMILES and graph. For SMILES representations each scaffold-molecule pair was transformed into two SMILES sequences which were then split into different tokens to denote atoms, bonds, or other tokens for grammar control (e.g. parentheses or numbers). All of these tokens were put together to form a vocabulary which recorded the index of each token (Fig. 1D). Here, we used the same conversion procedure and vocabulary as in *DrugEx v2*. In addition, we put a start token (GO) at the beginning of a batch of data as input and an end token (END) at the end of the same batch of data as output. After sequence padding with a blank token at empty positions, each SMILES sequence was rewritten as a series of token indices with a fixed length. Subsequently all of these sequences for both scaffolds and molecules were concatenated to construct the input and output matrix (Fig. 1E).

For the graph representation each molecule was represented as a five-row matrix, in which the first two rows stand for the index of the atom and bond types, respectively. The third and fourth rows represent the position of previous and current atoms connected by a bond (Fig. 1C). The columns of this matrix contain three sections to store scaffolds, growing parts, and linking parts. The scaffold section began with a start token in the first row and the last row was labelled the index of each scaffold starting from one. The fragments in the given scaffold for each molecule are put in the beginning of the matrix, followed by the growing part for the scaffold, and the last part is the connection bond between these

growing fragments with single bonds. For the growing part and linking sections, the last row was always zero and these two sections were separated by the column of end token. It is worth noticing that the last row was not directly involved in the training process. The vocabulary for graph representation was different from the SMILES format and it contains 38 atom types (Table S5.1) and four bond types (single, double, triple bonds and none). If the atom is the first occurrence in a given scaffold the type of the bond will be empty (indexed as 0 with token ‘*’). In addition, if the atom at the current position has been recorded in the matrix, the type of the atom will be empty. In order to grasp more details of the graph representation, we also provided the pseudocode for encoding (Table S5.2) and decoding (Table S5.3).

5.2.3. End-to-end deep learning

In this work, we compared three different sequential end-to-end DL architectures to deal with different molecular representations of either graph or SMILES (Fig. 5.2). These methods included: (A) Graph Transformer, (B) LSTM-based encoder-decoder model (LSTM-BASE), (C) LSTM-based encoder-decoder model with attention mechanisms (LSTM+ATTN) and (D) Sequential Transformer model. All of these DL models were constructed with *PyTorch* [27].

For SMILES representation three different models were constructed as follows (Fig. 5.2, right). The encoder and decoder in the LSTM-BASE model (Fig. 5.2B) had the same architectures, containing one embedding layer, three recurrent layers and one output layers (as we did for *DrugEx v2*) [24]. The number of neurons in the embedding and hidden layers were 128 and 512, respectively. The hidden states of the recurrent layer in the encoder are directly sent to the decoder as the initial states. On the basis of LSTM-BASE model, an attention layer was added between the encoder and decoder to form the LSTM+ATTN model (Fig. 5.2C). The attention layer calculates the weight for each position of the input sequence to determine which position the decoder needs to focus on during the decoding process. For each step, the weighted sums of the output calculated by the encoder are combined with the output of the embedding layer in the decoder to form the input for the

recurrent layers. The output of the recurrent layers is dealt with by the output layer to generate the probability distribution of tokens in the vocabulary in both of these two models.

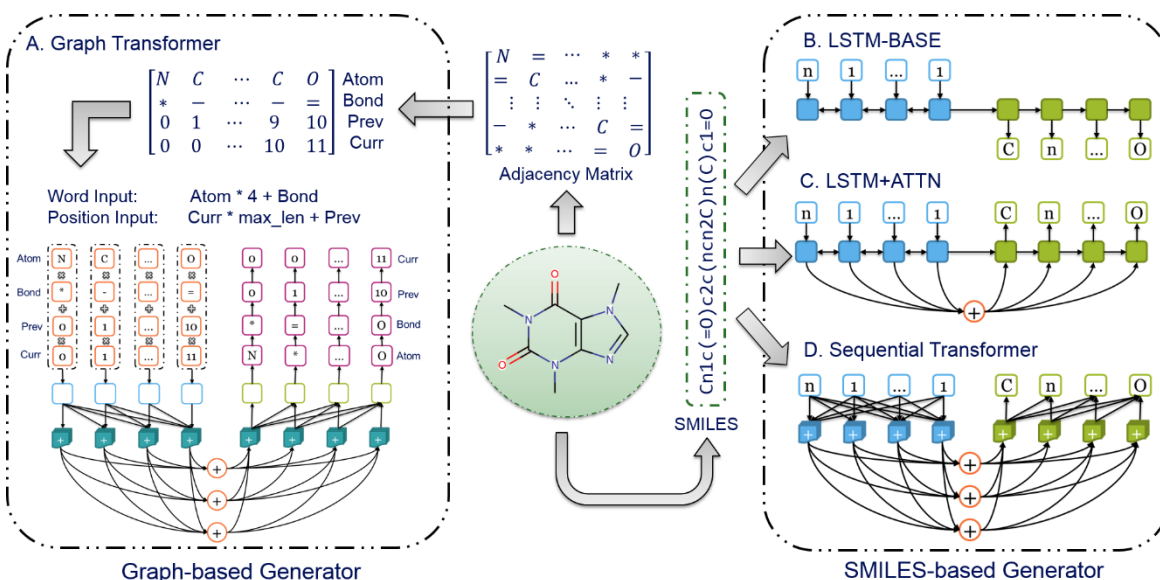


Fig. 5.2: Architectures of four different end-to-end deep learning models: (A) The Graph Transformer; (B) The LSTM-based encoder-decoder model (LSTM-BASE); (C) The LSTM-based encoder-decoder model with attention mechanisms (LSTM+ATTN); (D) The sequential Transformer model. The Graph Transformer accepts a graph representation as input and SMILES sequences are taken as input for the other three models.

The sequential Transformer has a distinct architecture compared to the LSTM+ATTN model although it also exploits an attention mechanism. For the embedding layers “position encodings” are added into the typical embedding structure as the first layer of the encoder and decoder. This ensures that the model no longer needs to encode the input sequence token by token but can process all tokens in parallel. For the position embedding, sine and cosine functions are used to define its formula as follows:

$$PE_{(p,2i)} = \sin(pos/10000^{2i/d_m})$$

$$PE_{(p,2i+1)} = \cos(pos/10000^{2i/d_m})$$

where $PE(p, i)$ is the i^{th} dimension of the position encoding at position p . It has the same dimension $d_m = 512$ as the typical embedding vectors so that the two can be summed.

In addition, self-attention is used in the hidden layers to cope with long-range dependencies.

For each hidden layer in the encoder, it employs a residual connection around a multi-head self-attention sublayer and feed-forward sublayer followed by layer normalization. Besides these two sublayers in the decoder a third sublayer with multi-head attention is inserted to capture the information from output of the encoder.

This self-attention mechanism is defined as the scaled dot-product attention with three vectors: queries (Q), keys (K) and values (V), of which the dimensions are d_q , d_k , d_v , respectively. The output matrix is computed as:

$$Attention(Q, K, V) = softmax\left(\frac{QK^T}{\sqrt{d_k}}\right)V$$

Instead of a single attention function, the Transformer adopts multi-head attention to combine information from different representations at different positions which is defined as:

$$MultiHead(Q, K, V) = Concat(head_1, \dots, head_h)W^O$$

where h is the number of heads. For each head, the attention values were calculated by different and learned linear projections with Q , K and V as follows:

$$head_i = Attention(QW_i^Q, KW_i^K, VW_i^V)$$

where W^O , W^Q , W^K and W^V are metrics of learned weights and we set $h = 8$ as the number of heads and $d_k = d_v = 64$ in this work.

For the graph representation of molecules, we updated the sequential Transformer structure to propose a Graph Transformer (Fig. 5.2A). Similar to the sequential Transformer the Graph Transformer also requires the encodings of both word and position as the input. For the input word, the atom and bond cannot be processed simultaneously; therefore we combined the index of atom and bond together and defined it as follows:

$$I = I_{atom} \times 4 + I_{bond}$$

meaning the index of the input word (I) calculating word vectors are calculated from atom index (I_{atom}) multiplied by four (the total number of bond types defined) and add the bond index (I_{bond}). Similarly, the position of each step cannot be used to calculate the position encoding directly. Faced with more complex data structure than sequential data,

Dosovitskiy *et al.* proposed a new positional encoding scheme to define the position for each patch in image data for image recognition [28]. Inspired by their work the position encoding at each step was defined as:

$$P = P_{curr} \times L_{max} + P_{prev}$$

here the input position (P) for calculating the position encoding was the current position (P_{curr}) multiplied by the max length (L_{max}) and adding the previous position (P_{prev}), which was then processed with the same positional encoding method as with the sequential Transformer. For the decoder, the hidden vector from the transformer was taken as the starting point to be decoded by a GRU-based recurrent layer; and the probability of atom, bond, previous and current position was decoded one by one sequentially.

When graph-based molecules are generated, the chemical valence rule is checked in every step. The invalid values of atom and bond types will be masked and an incorrect previous and current position will be removed ensuring the validity of all generated molecules. It is worth noticing that before being encoded, each molecule will be kekulized, meaning that the aromatic rings will be inferred to transform into either single or double bonds. The reason for this is that aromatic bonds interfere with the calculation of the valence value for each atom.

During the training process of SMILES-based models, the negative log likelihood function was used to construct the loss function to guarantee that the token in the output sequence had the largest probability to be chosen. In comparison, the loss function used by the Graph Transformer model also contains four parts for atom, bond, previous and current sites. And the sum of these negative log probability values is minimized to optimize the parameters in the model. For this, the Adam algorithm was used for the optimization of the loss function. Here, the learning rate was set at 10^{-4} , the batch size was 256, and training steps were set to 20 epochs for pre-training and 1,000 epochs for fine-tuning.

5.2.4. Multi-objective optimization

In order to combine multiple objectives we exploited a Pareto-based ranking algorithm

with GPU acceleration as mentioned in *DrugEx* v2 [24]. Given two solutions m_1 and m_2 with their scores $(x_1, x_2, \dots, x_n)$ and $(y_1, y_2, \dots, y_n)$, then m_1 is said to Pareto dominate m_2 if and only if:

$$\forall j \in \{1, \dots, n\}: x_j \geq y_j \text{ and } \exists j \in \{1, \dots, n\}: x_j > y_j$$

otherwise, m_1 and m_2 are non-dominated with each other. After the dominance between all pair of solutions being determined, the non-dominated scoring algorithm is exploited to obtain a rank of Pareto frontiers which consist of a set of solutions. After obtaining frontiers between dominant solutions molecules were ranked based on the average Tanimoto-distance with other molecules instead of crowding distance in the same frontier. Subsequently molecules with smaller distances were ranked on the top. The final reward R^* is defined as:

$$R^* = \begin{cases} 0.5 + \frac{k - N_{undesired}}{2N_{desired}}, & \text{if desired} \\ \frac{k}{2N_{undesired}}, & \text{if undesired} \end{cases}$$

here k is the index of the solution in the Pareto rank and rewards of undesired and desired solutions will be evenly distributed in $(0, 0.5]$ and $(0.5, 0.1]$, respectively.

In this work, we took two objectives into consideration: 1) QED score [29] as implemented by RDKit (from 0 to 1) to evaluate the drug-likeness of each molecule (a larger value means more drug-like) ; 2) an affinity score towards A_{2A}AR which was implemented by a random forest regression model with Scikit-Learn [30] like in *DrugEx* v2. The input descriptors consisted of 2048D ECFP6 fingerprints and 19D physico-chemical descriptors (PhysChem). PhysChem included: molecular weight, logP, number of H bond acceptors and donors, number of rotatable bonds, number of amide bonds, number of bridge head atoms, number of hetero atoms, number of spiro atoms, number of heavy atoms, the fraction of SP3 hybridized carbon atoms, number of aliphatic rings, number of saturated rings, number of total rings, number of aromatic rings, number of heterocycles, number of valence electrons, polar surface area, and Wildman-Crippen MR value. Again it was determined if generated molecules are desired based on the Affinity score (larger than the threshold = 6.5). In addition, the SA score was also exploited to evaluate the synthesizability of generated molecules, which is also calculated by RDKit [31].

5.2.5. Reinforcement learning

In this work we constructed a reinforcement learning framework based on the interplay between the Graph Transformer (agent) and the two scoring functions (environment). A policy gradient method was implemented to train the reinforcement learning model, the objective function is designated as follows:

$$J(\theta) = \mathbb{E}[R^*(y_{1:T})|\theta] = \sum_{t=1}^T \log G(y_t|y_{1:t-1}) \cdot R^*(y_{1:T})$$

here for each step t during the generation process, the generator (G) determines the probability of each token (y_t) from the vocabulary to be chosen based on the generated sequence in previous steps ($y_{1:t-1}$). In the sequence-based models y_t can only be a token in the vocabulary to construct SMILES while it can be different type of atoms or bonds or the previous or current position in the graph-based model. The parameters in its objective function are updated by employing a policy gradient based on the expected end reward (R^*) received from the predictor. By maximizing this function the parameter θ in the generator can be optimized to ensure that the generator designs desired molecules which obtain a high reward score.

In order to improve the diversity and reliability of generated molecules, we implemented our exploration strategy for molecule generation during the training loops. In the training loop our generator is trained to produce the chemical space as defined by the target of interest. In this strategy there are two networks with the same architectures, an exploitation net (G_θ) and an exploration net (G_φ). G_φ did not need to be trained and its parameters were always fixed and it is based on the general drug-like chemical space for diverse targets obtained from ChEMBL. The parameters in G_θ on the other hand were updated for each epoch based on the policy gradient. Again an *exploring rate* (ε) was defined with a range of [0.0, 1.0] to determine the percentage of scaffolds being randomly selected as input by G_φ to generate molecules. Conversely G_θ generated molecules with other input scaffolds. After the training process was finished G_φ was removed and only G_θ was left as the final model for molecule generation.

5.2.6. Performance evaluation

In order to evaluate the performance of the generators, four coefficients were calculated from the population of generated molecules (validity, accuracy, desirability, and uniqueness) which are defined as:

$$\text{Validity} = \frac{N_{\text{valid}}}{N_{\text{total}}}$$

$$\text{Accuracy} = \frac{N_{\text{accurate}}}{N_{\text{total}}}$$

$$\text{Desirability} = \frac{N_{\text{desired}}}{N_{\text{total}}}$$

$$\text{Uniqueness} = \frac{N_{\text{unique}}}{N_{\text{total}}}$$

here N_{total} is the total number of molecules, N_{valid} is the number of molecules parsed as valid SMILES sequences, N_{accurate} is the number of molecules that contained given scaffolds, N_{desired} is the number of desired molecules that reach all required objectives, and N_{unique} is the number of molecules which are different from others in the dataset .

To measure molecular diversity, we adopted the Solow Polasky measurement as in *DrugEx* v2 [24]. This approach was proposed by Solow and Polasky in 1994 to estimate the diversity of a biological population in an eco-system [32]. The formula to calculate diversity was redefined to normalize the range of values from $[1, m]$ to $(0, m]$ as follows:

$$I(A) = \frac{1}{|A|} \mathbf{e}^T F(\mathbf{s})^{-1} \mathbf{e}$$

where A is a set of drug molecules with a size of $|A|$ equal to m , $\mathbf{e}$ is an m -vector of 1's and $F(\mathbf{s}) = [f(d_{ij})]$ is a non-singular $m \times m$ distance matrix, in which $f(d_{ij})$ stands for the distance function of each pair of molecule provided as follows:

$$f(d) = e^{-\theta d_{ij}}$$

here we defined the distance d_{ij} of molecules s_i and s_j by using the Tanimoto-distance with ECFP6 fingerprints as follows:

$$d_{ij} = d(s_i, s_j) = 1 - \frac{|s_i \cap s_j|}{|s_i \cup s_j|},$$

where $|s_i \cap s_j|$ represents the number of common fingerprint bits, and $|s_i \cup s_j|$ is the number of union fingerprint bits.

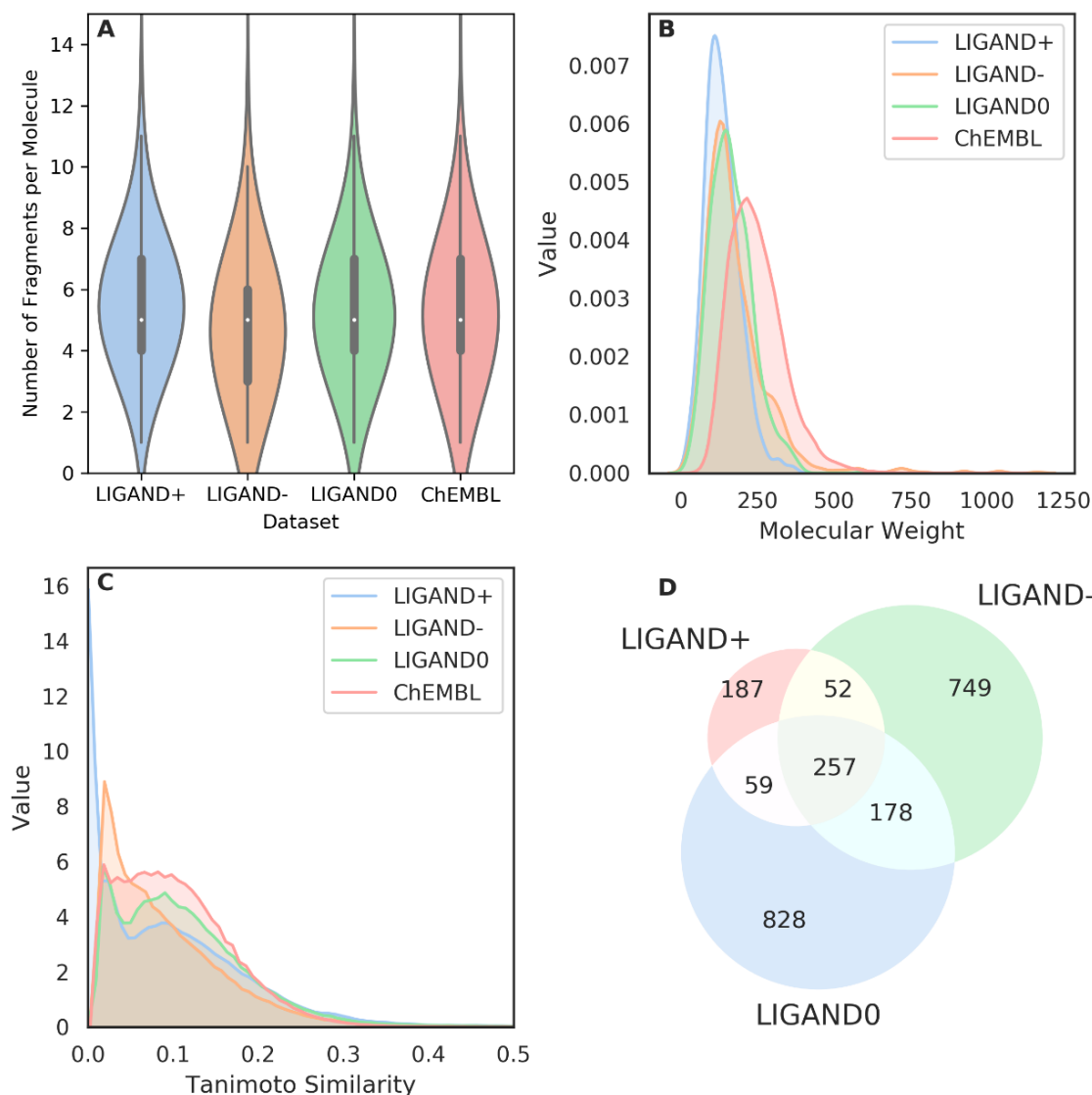


Fig 5.3: Analysis of some properties of fragments in the *ChEMBL* set and three *LIGAND* subsets. (A) Violin plot for the distribution of the number of fragments per molecules; (B) Distribution of molecular weight of these fragments; (C) Distribution of the similarity of the fragments measured by the Tanimoto-similarity with ECFP4 fingerprints; (D) Venn diagram for the intersection of the fragments existing in the three subsets of the *LIGAND* set.

5.3. Results and discussion

5.3.1. Fragmentation of molecule

As stated we decomposed each molecule into a series of fragments with the BRICS algorithm to construct scaffold-molecule pairs. Within BRICS each organic compound can

be split into retrosynthetically interesting chemical substructures with a compiled elaborate set of rules. For the *ChEMBL* and *LIGAND* sets, we respectively obtained 194,782 and 2,223 fragments. We further split the *LIGAND* set into three parts: active ligands (*LIGAND*⁺, 2,638), inactive ligands (*LIGAND*⁻, 2710) and undetermined ligands (*LIGAND*⁰, 5480) based on the pX of bioactivity for A_{2A}AR. The number of fragments in these four datasets have a similar distribution and there are approximately five fragments on average for each molecule with a 95% confidence between [0, 11] (Fig. 5.3A).

In the *LIGAND* set the three subsets have a similar molecular weight distribution of the fragments (Fig. 5.3B) while the average is 164.3Da, smaller than in the *ChEMBL* set (247.3Da). In order to check the similarity of these fragments we used the Tanimoto similarity calculation with ECFP4 fingerprints between each pair of fragments in the same dataset. We found that most of them were smaller than 0.5 indicating that they are dissimilar to each other (Fig. 5.3C). Especially, the fragments in the *LIGAND*⁺ set have the largest diversity. Moreover, the distribution of different fragments in these three subsets of the *LIGAND* set are shown in Fig. 5.3D. The molecules in these three subsets have their unique fragments and share some common substructures.

5.3.2. Pre-training & fine-tuning

After finishing the dataset construction, four models were pre-trained on the *ChEMBL* set and fine-tuned on the *LIGAND* set. Here, these models were benchmarked on a server with four GTX1080Ti GPUs. After the training process converged each fragment in the test set was presented as input for 10 times to generate molecules. The performance is shown in Table 5.1. The training of Transformer models was faster but consumed more computational resources than LSTM-based methods. In addition, Transformer methods outperformed LSTM-based methods using SMILES. Although the three SMILES-based models improved after being fine-tuned they were still outperformed by the Graph Transformer because of the advantages of the graph representation. To further check the accuracy of generated molecules we also compared the chemical space between the generated molecules and the compounds in the training set with three different

representations 1) MW $\sim$ logP; 2) PCA with 19D PhysChem descriptors; 3) tSNE with 2048D ECFP6 fingerprints (Fig. 5.4). The region occupied by molecules generated by the Graph Transformer overlapped completely with the compounds in both the *ChEMBL* and *LIGAND* sets.

Table 5.1: The performance of four different generators for pre-training and fine-tuning processes.

Methods	Pre-trained Model		Fine-tuned Model		Time	Memory
	Validity	Accuracy	Validity	Accuracy		
Graph Transformer	100%	99.3%	100%	99.2%	453.8 s	14.5 GB
Sequential Transformer	96.7%	72.0%	99.3%	95.7%	832.3 s	31.7 GB
LSTM-BASE	93.9%	44.1%	98.7%	91.8%	834.6 s	5.5 GB
LSTM+ATTN	89.7%	52.2%	96.4%	90.2%	1212.5 s	15.9 GB

The graph representation for molecules has more advantages over the SMILES representation when dealing with fragment-based molecule design: 1) **Invariance in the local scale**: During the process of molecule generation multiple fragment in the given scaffold can be put into any position in the output matrix without changing the order of atoms and bonds in that fragment. 2) **Extendibility in the global scale**: When the fragments in the scaffold are growing or being linked, they can be flexibly appended in the end column of the graph matrix while the original data structure does not need changing. 3) **Free of grammar**: Unlike in SMILES sequences there is no explicit grammar to constrain the generation of molecules, such as the parentheses for branches and the numbers for rings in SMILES; 4) **Accessibility of chemical rules**: For each added atom or bond the algorithm can detect if the valence of atoms is valid or not and mask invalid atoms or bonds in the vocabulary to guarantee the whole generated matrix can be successfully parsed into a molecule. With these advantages the Graph Transformer generates molecules faster while using less memory.

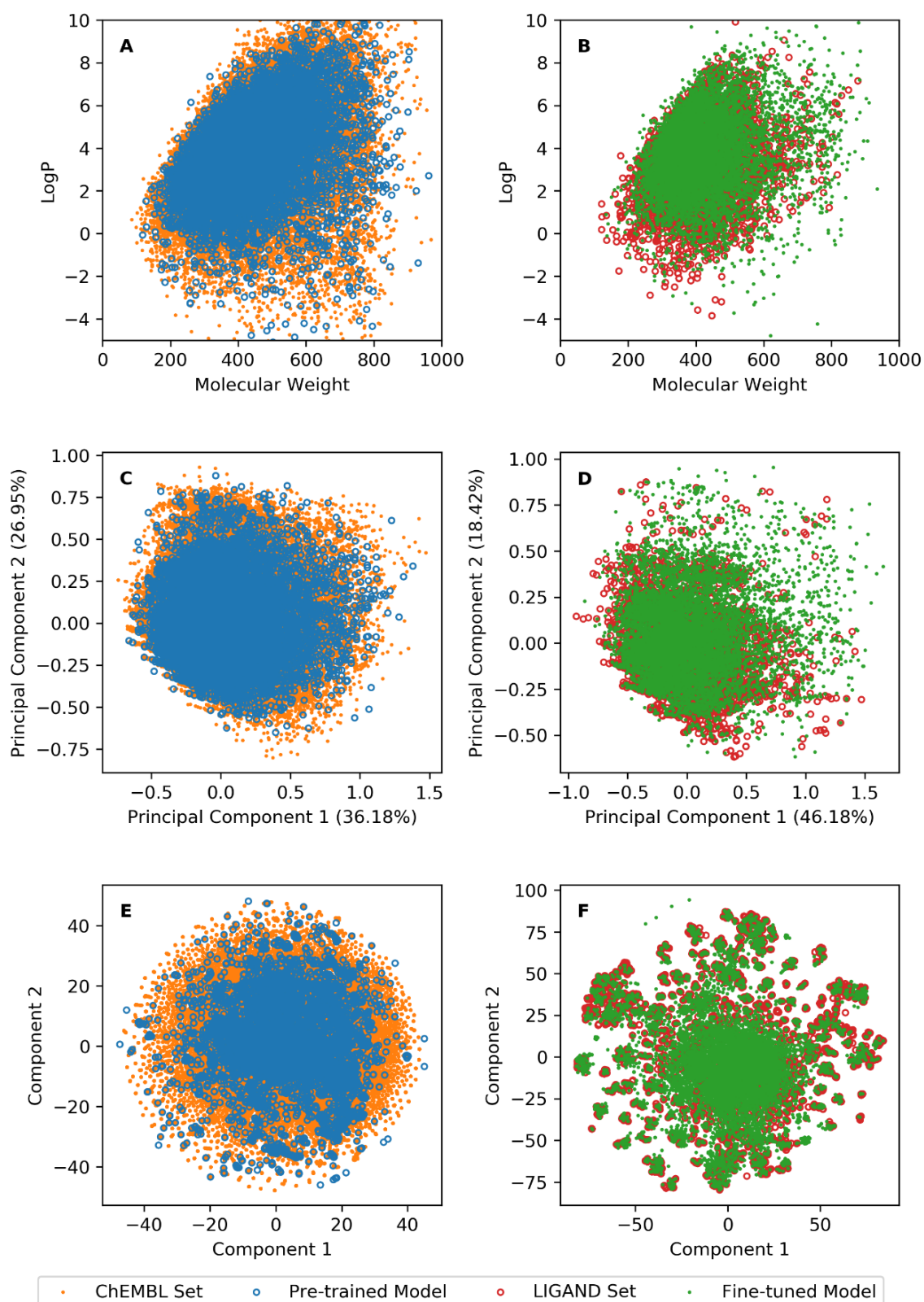


Fig. 5.4: The chemical space of generated molecules by the Graph Transformer pre-trained on the ChEMBL set (A, C and E) and being fine-tuned on the LIGAND set (B, D and F). Chemical space was represented by either logP ~ MW (A, B) and first two components in PCA on PhysChem descriptors (C, D) and t-SNE on ECFP6 fingerprints (E, F).

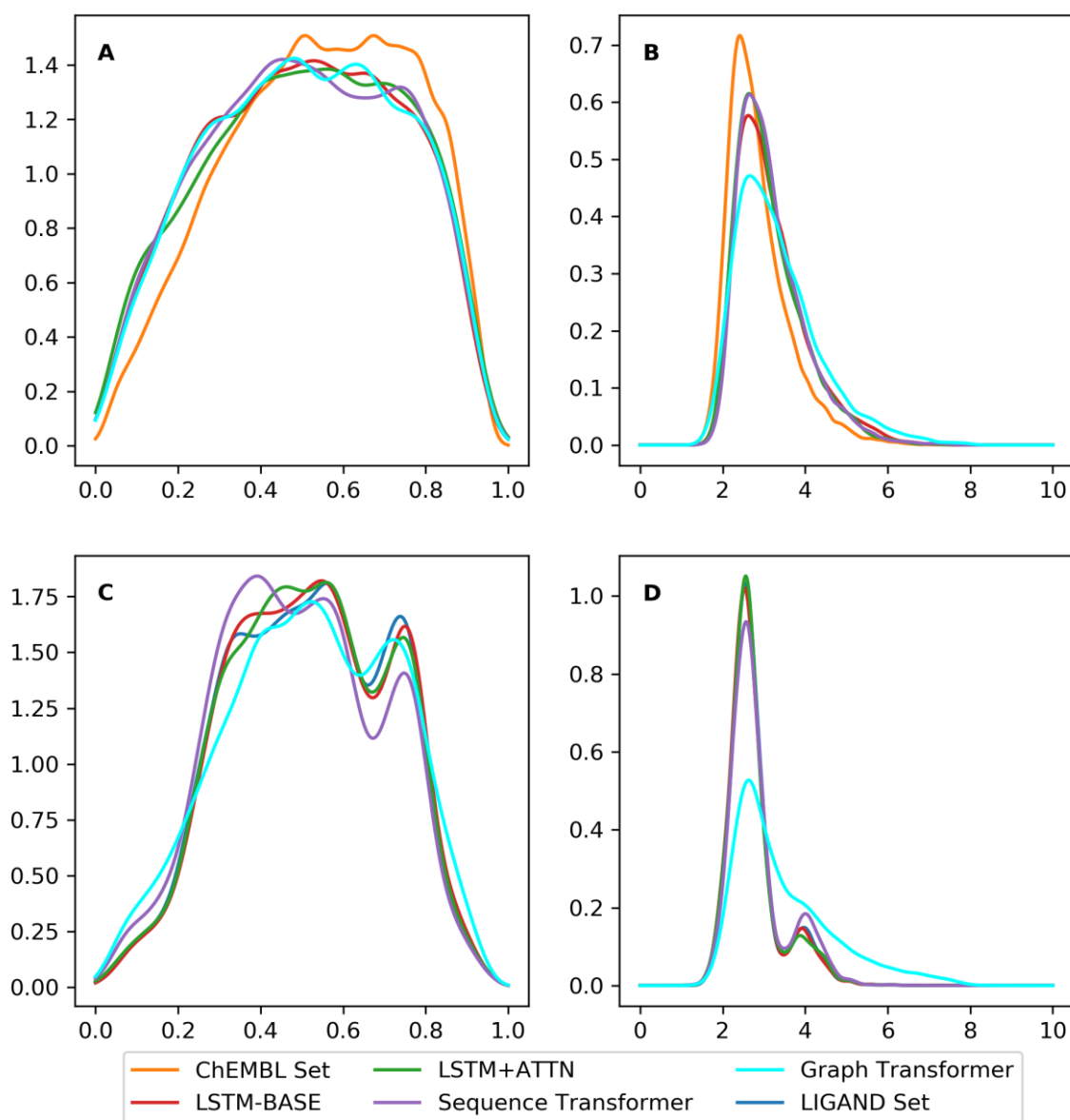


Fig. 5.5: the distribution of QED score (A, C) and SA score (B, D) of desired ligands in the *LIGAND* set and of molecules generated by four different generators.

However, after examining the QED scores and SA scores, we found that although the distribution of QED scores was similar to each other, the synthesizability of the molecules generated by the Graph Transformer were no better than the SMILES-based generators, especially when fine-tuning on the *LIGAND* set (Fig. 5.5). The possible reason is that the molecules generated by the Graph Transformer contains some uncommon rings when the model dealt with long-distance dependencies. In addition, because of more complicated data structure and more parameters in the model, the synthesizability performance of Graph Transformer was not considered high enough when being trained on the small dataset (e.g. the *LIGAND* set). It is also worth noticing that there still was a small fraction of generated

molecules that did not contain the given scaffolds. This is caused by the kekulization problem. For example, a scaffold ‘CCC’ can be grown into ‘C1=C(C)C=CC=C1’. After being sanitized, it can be transformed into ‘c1c(C)cccc1’. In this process one single bond in the scaffold is changed to an aromatic bond, which causes the mismatch between the scaffold and the molecule. Currently our algorithm cannot solve this problem because if the aromatic bond is taken into consideration, the valence of aromatic atoms is difficult to be calculated accurately. This would lead to the generation of invalid molecules. Therefore, there is no aromatic bond provided in the vocabulary and all of the aromatic rings are inferred automatically through the molecule sanitization method in RDKit.

5.3.3. Policy gradient

Because the Graph Transformer generates molecules accurately and fast it was chosen as the agent in the RL framework. Two objectives were tested in the training process of this work. The first one was affinity towards A_{2A}AR, which is predicted by the random forest-based regression model from *DrugEx* v2; the second one was the QED score calculated with RDKit to measure how similar the generated molecule is to known approved drugs. With the policy gradient method as the reinforcement learning framework two cases were tested. On the one hand, predicted affinity for A_{2A}AR was considered without the QED score. On the other hand, both objectives were used to optimize the model with Pareto ranking. In the first case 86.1% of the generated molecules were predicted active, while the percentage of predicted active molecules in the second case was 74.6%. Although the generator generated more active ligands without the QED score constraint most of them are not drug-like as they always have a molecular weight larger than 500Da. However, when we checked the chemical space represented by tSNE with ECFP6 fingerprints the overlap region between generated molecules and ligands in the training set was not complete implying that they fall out of the applicability domain of the regression model.

In the version of v2, we provided an exploration strategy which simulated the idea of evolutionary algorithms such as *crossover* and *mutation* manipulations [24]. However, when coupled to the Graph Transformer there were some difficulties and we had to give up this strategy. Firstly, the mutation strategy did not improve with different mutation rates. A

possible reason is that before being generated part the molecule was fixed with a given scaffold, counteracting the effect of mutation caused by the mutation net. Secondly, the *crossover* strategy is computationally very expensive in this context. This strategy needs the convergence of model training and iteratively updates the parameters in the agent. With multiple iterations, it takes a long period of time beyond the computational resources we can currently access. As a result, we updated the exploration strategy as mentioned in the Methods section with six different exploration rates: [0.0, 0.1, 0.2, 0.3, 0.4, 0.5].

Table 5.2: the performance of the Graph Transformer with different exploration rates in the RL framework.

ϵ	Accuracy	Desirability	Uniqueness	Diversity
0.0	99.7%	74.6%	60.7%	0.879
0.1	99.7%	66.8%	75.0%	0.842
0.2	99.8%	61.6%	80.2%	0.879
0.3	99.7%	56.8%	89.8%	0.874
0.4	99.7%	54.8%	88.8%	0.859
0.5	99.7%	46.8%	88.5%	0.875

Changes to the exploration rate do not influence accuracy and have a low effect on diversity. However desirability (finding active ligands) and uniqueness can be influenced significantly. Empirically determining an optimal value for a given chemical space is recommended.

After training of the models, the scaffolds in the test set were input 10 times to generate molecules. The results for accuracy, desirability, uniqueness, and diversity with different exploration rates are shown in Table 5.2. With a low ϵ the model generates more desired molecules, but the uniqueness of the generated molecules can be improved. At $\epsilon = 0.3$ the model generated the highest percentage of unique desired molecules (56.8%). Diversity was always larger than 0.84 and the model achieved the largest value (0.88) with $\epsilon = 0.0$ or $\epsilon = 0.2$. The chemical space represented by tSNE with ECFP6 fingerprints confirms that our exploration strategy produces a set of generated molecules completely covering the region occupied by the *LIGAND* set (Fig. 5.6).

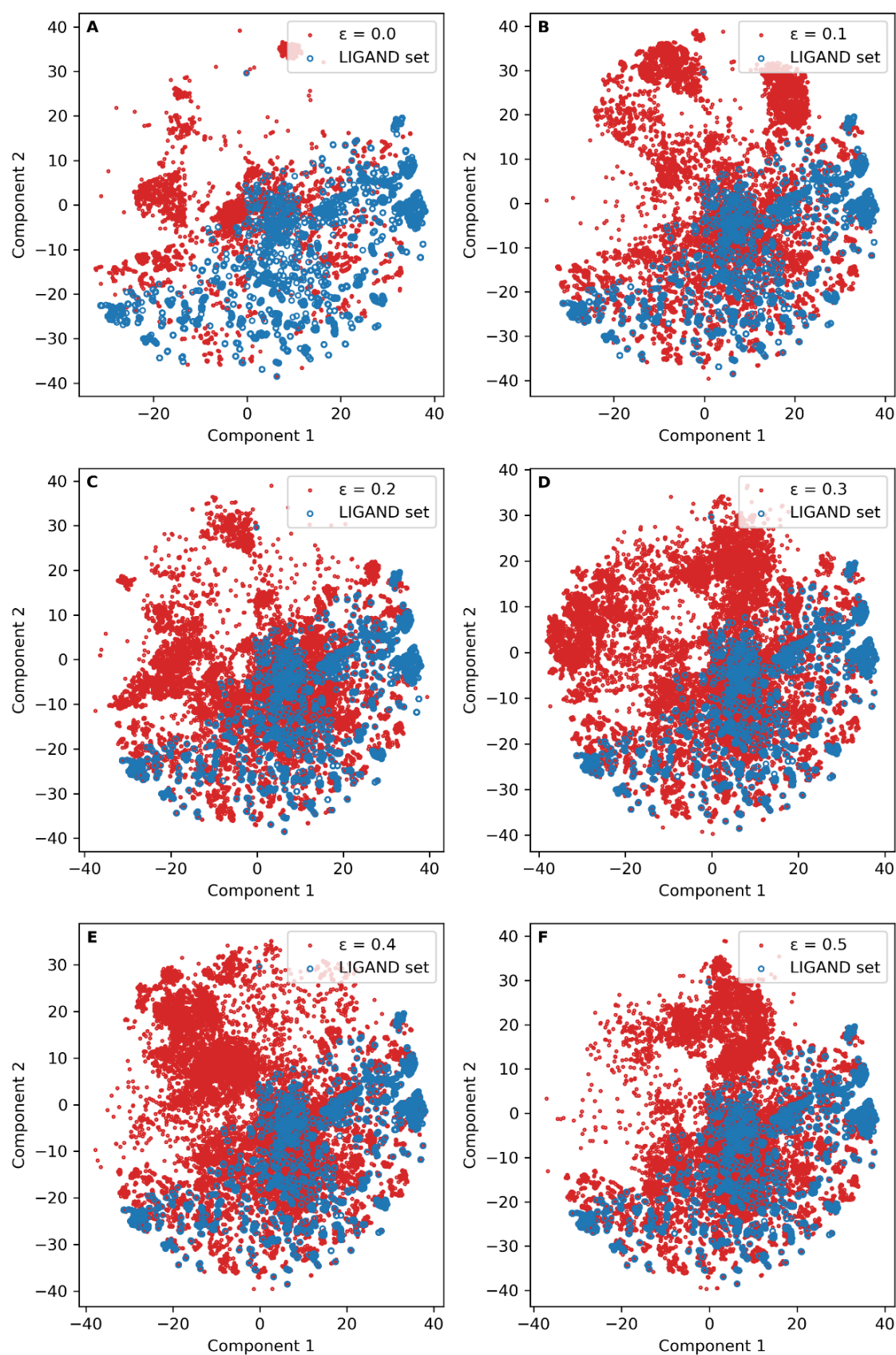


Fig. 5.6: The chemical space of generated molecules by the Graph Transformer trained with different exploration rates in the RL framework. The chemical space was represented by t-SNE on ECFP6 fingerprints.

5.3.4. Generated molecules

In the chemical space for antagonists of $A_{2A}AR$, furan, triazine, aminotriazole, and purine derivatives such as xanthine and azapurine are common fragments. The Graph Transformer model produced active ligands for $A_{2A}AR$ (inferred from the predictors) with different combinations of these fragments as the scaffolds. Taking these molecules generated by the Graph Transformer as an example, we filtered out the molecules with potentially reactive groups (such as aldehydes) and uncommon ring systems and listed 30 desired molecules as putative $A_{2A}AR$ ligands/antagonists (Fig. 5.7). For each scaffold, five molecules were selected and assigned in the same row. These molecules are considered a valid starting point for further considerations and work (e.g. molecular docking or simulation).

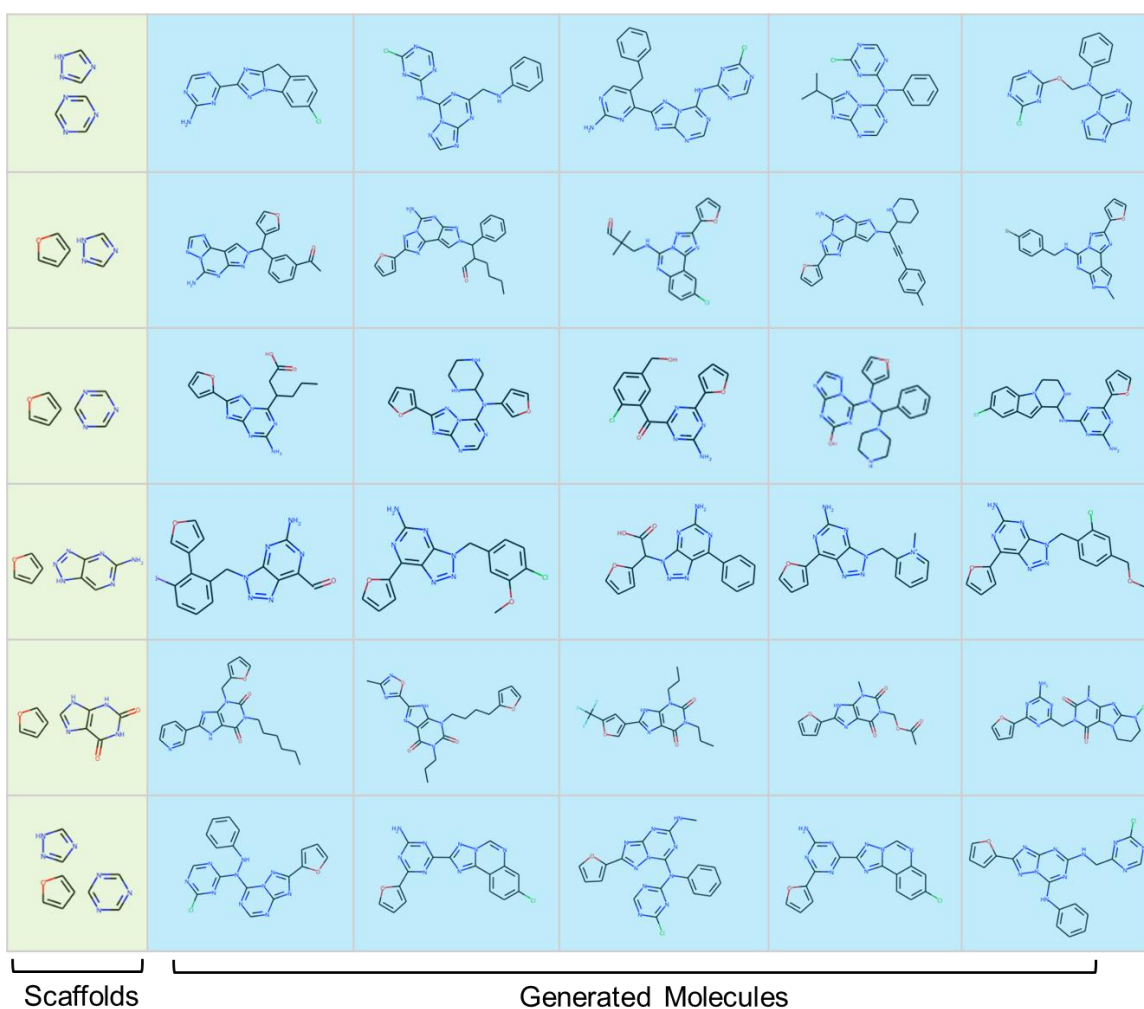


Fig. 5.7: Sample of generated molecules with the Graph Transformer with different scaffolds. These scaffolds include: furan, triazine, aminotriazole, xanthine and azapurine. The generated molecules based on the same scaffolds are aligned in the same row.

5.4. Conclusion and Future Perspective

In this study, *DrugEx* was updated with the ability to design novel molecules based on the scaffolds containing multiple fragments as input. In this version (v3), a new positional encoding scheme for atoms and bonds was proposed to make the Transformer model deal with a molecular graph representation. With one model multiple fragments in the scaffold can be grown at the same time and connected to generate a new molecule. In addition, chemical rules on valence are enforced at each step of the process of molecule generation to ensure that all generated molecules are valid. This is impossible for SMILES-based generation, as SMILES-based molecules are constrained by grammar that allows a 2D topology to be represented in a sequential way. With multi-objective reinforcement learning the model generates drug-like ligands, in our case for the A_{2A}AR target.

In future work, the Graph Transformer will be extended to include other information as input to design drugs conditionally. For example, proteochemometric modelling (PCM) can take information for both ligands and targets as input to predict the affinity of their interactions, which allows promiscuous (useful for e.g., viral mutants) or selective (useful for e.g., kinase inhibitors) properties [33]. The Transformer can then be used to construct inverse PCM models which take the protein information as input (e.g. sequences, structures or descriptors) to design active ligands for a given protein target without known ligands. Moreover, the Transformer can also be used for lead optimization. For instance, the input can be a “hit” already, generating “optimized” ligands, or a “lead” with side effects to produce ligands with a better ADME/tox profile.

Declarations

Availability of data and materials

The data used in this study is publicly available ChEMBL data, the algorithm published in this manuscript is made available at <https://github.com/XuhanLiu/DrugEx>.

Authors' Contributions

XL and GJPvW conceived the study and performed the experimental work and analysis. KY, APIJ and HWTvV provided feedback and critical input. All authors read, commented on and approved the final manuscript.

Acknowledgements

XL thanks Chinese Scholarship Council (CSC) for funding, GJPvW thanks the Dutch Research Council and Stichting Technologie Wetenschappen (STW) for financial support (STW-Veni #14410). Thanks go to Dr. Xue Yang for verifying Table S1 and Dr. Anthe Janssen checking the convergence of tSNE. We also acknowledge Bert Beerkens for providing the common scaffolds used to generate molecules as an example.

Competing Interests

The authors declare that they have no competing interests

References

- Polishchuk PG, Madzhidov TI, Varnek A (2013) Estimation of the size of drug-like chemical space based on GDB-17 data. *J Comput Aided Mol Des* 27 (8):675-679. doi:10.1007/s10822-013-9672-4
- Hajduk PJ, Greer J (2007) A decade of fragment-based drug design: strategic advances and lessons learned. *Nat Rev Drug Discov* 6 (3):211-219. doi:10.1038/nrd2220
- Card GL, Blasdel L, England BP, Zhang C, Suzuki Y, Gillette S, Fong D, Ibrahim PN, Artis DR, Bollag G, Milburn MV, Kim SH, Schlessinger J, Zhang KY (2005) A family of phosphodiesterase inhibitors discovered by cocrystallography and scaffold-based drug design. *Nat Biotechnol* 23 (2):201-207. doi:10.1038/nbt1059
- Bian Y, Xie XS (2018) Computational Fragment-Based Drug Design: Current Trends, Strategies, and Applications. *AAPS J* 20 (3):59. doi:10.1208/s12248-018-0216-7
- Hughes JP, Rees S, Kalindjian SB, Philpott KL (2011) Principles of early drug discovery. *Br J Pharmacol* 162 (6):1239-1249. doi:10.1111/j.1476-5381.2010.01127.x
- Fredholm BB (2010) Adenosine receptors as drug targets. *Exp Cell Res* 316 (8):1284-1288. doi:10.1016/j.yexcr.2010.02.004
- Chen JF, Eltzschig HK, Fredholm BB (2013) Adenosine receptors as drug targets--what are the challenges? *Nat Rev Drug Discov* 12 (4):265-286. doi:10.1038/nrd3955
- Moro S, Gao ZG, Jacobson KA, Spalluto G (2006) Progress in the pursuit of therapeutic adenosine receptor antagonists. *Med Res Rev* 26 (2):131-159. doi:10.1002/med.20048
- Jespers W, Oliveira A, Prieto-Diaz R, Majellaro M, Aqvist J, Sotelo E, Gutierrez-de-Teran H (2017) Structure-Based Design of Potent and Selective Ligands at the Four Adenosine Receptors. *Molecules* 22 (11). doi:10.3390/molecules22111945
- Sheng C, Zhang W (2013) Fragment informatics and computational fragment-based drug design: an overview and update. *Med Res Rev* 33 (3):554-598. doi:10.1002/med.21255
- LeCun Y, Bengio Y, Hinton G (2015) Deep learning. *Nature* 521 (7553):436-444. doi:10.1038/nature14539
- Liu X, IJzerman AP, van Westen GJP (2021) Computational Approaches for De Novo Drug Design: Past, Present, and Future. *Methods Mol Biol* 2190:139-165. doi:10.1007/978-1-0716-0826-5_6
- Gomez-Bombarelli R, Wei JN, Duvenaud D, Hernandez-Lobato JM, Sanchez-Lengeling B, Sheberla D, Aguilera-Iparraguirre J, Hirzel TD, Adams RP, Aspuru-Guzik A (2018) Automatic Chemical Design Using a Data-Driven Continuous Representation of Molecules. *ACS Cent Sci* 4 (2):268-276. doi:10.1021/acscentsci.7b00572
- Segler MHS, Kogej T, Tyrchan C, Waller MP (2018) Generating Focused Molecule Libraries for Drug Discovery with Recurrent Neural Networks. *ACS Cent Sci* 4 (1):120-131. doi:10.1021/acscentsci.7b00512
- Benjamin S-L, Carlos O, Gabriel L. G, Alan A-G (2017) Optimizing distributions over molecular space. An Objective-Reinforced Generative Adversarial Network for Inverse-design Chemistry (ORGANIC). doi:10.26434/chemrxiv.5309668.v3
- Olivecrona M, Blaschke T, Engkvist O, Chen H (2017) Molecular de-novo design through deep reinforcement learning. *Journal of cheminformatics* 9 (1):48. doi:10.1186/s13321-017-0235-x
- Blaschke T, Arus-Pous J, Chen H, Margreitter C, Tyrchan C, Engkvist O, Papadopoulos K, Patronov A (2020) REINVENT 2.0: An AI Tool for De Novo Drug Design. *Journal of chemical information and modeling* 60 (12):5918-5922. doi:10.1021/acs.jcim.0c00915
- Lim J, Hwang SY, Moon S, Kim S, Kim WY (2019) Scaffold-based molecular design with a graph generative model. *Chem Sci* 11 (4):1153-1164. doi:10.1039/c9sc04503a

19. Li Y, Hu J, Wang Y, Zhou J, Zhang L, Liu Z (2020) DeepScaffold: A Comprehensive Tool for Scaffold-Based De Novo Drug Discovery Using Deep Learning. *Journal of chemical information and modeling* 60 (1):77-91. doi:10.1021/acs.jcim.9b00727
20. Arus-Pous J, Patronov A, Bjerrum EJ, Tyrchan C, Reymond JL, Chen H, Engkvist O (2020) SMILES-based deep generative scaffold decorator for de-novo drug design. *Journal of cheminformatics* 12 (1):38. doi:10.1186/s13321-020-00441-8
21. Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, Kaiser L, Polosukhin IJae-p (2017) Attention Is All You Need. arXiv:1706.03762
22. Yang Y, Zheng S, Su S, Zhao C, Xu J, Chen H (2020) SyntaLinker: automatic fragment linking with deep conditional transformer neural networks. *Chem Sci* 11 (31):8312-8322. doi:10.1039/d0sc03126g
23. Liu X, Ye K, van Vlijmen HWT, IJzerman AP, van Westen GJP (2019) An exploration strategy improves the diversity of de novo ligands using deep reinforcement learning: a case for the adenosine A2A receptor. *Journal of cheminformatics* 11 (1):35. doi:10.1186/s13321-019-0355-6
24. Liu X, Ye K, van Vlijmen HWT, Emmerich M, IJzerman AP, van Westen GJP (2021) DrugEx v2: de novo design of drug molecules by Pareto-based multi-objective reinforcement learning in polypharmacology. *Journal of cheminformatics* 13 (1):85. doi: 10.1186/s13321-021-00561-9
25. Gaulton A, Bellis LJ, Bento AP, Chambers J, Davies M, Hersey A, Light Y, McGlinchey S, Michalovich D, Al-Lazikani B, Overington JP (2012) ChEMBL: a large-scale bioactivity database for drug discovery. *Nucleic Acids Res* 40 (Database issue):D1100-1107. doi:10.1093/nar/gkr777
26. Degen J, Wegscheid-Gerlach C, Zaliani A, Rarey M (2008) On the art of compiling and using 'drug-like' chemical fragment spaces. *ChemMedChem* 3 (10):1503-1507. doi:10.1002/cmdc.200800178
27. PyTorch. <https://pytorch.org/>.
28. Dosovitskiy A, Beyer L, Kolesnikov A, Weissenborn D, Zhai X, Unterthiner T, Dehghani M, Minderer M, Heigold G, Gelly S, Uszkoreit J, Houlsby NJae-p (2020) An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale. arXiv:2010.11929
29. Bickerton GR, Paolini GV, Besnard J, Muresan S, Hopkins AL (2012) Quantifying the chemical beauty of drugs. *Nat Chem* 4 (2):90-98. doi:10.1038/nchem.1243
30. Scikit-Learn: machine learning in Python. <http://www.scikit-learn.org/>.
31. Ertl P, Schuffenhauer A (2009) Estimation of synthetic accessibility score of drug-like molecules based on molecular complexity and fragment contributions. *Journal of cheminformatics* 1 (1):8. doi:10.1186/1758-2946-1-8
32. Solow AR, Polasky S (1994) Measuring biological diversity. *Environmental and Ecological Statistics* 1 (2):95-103. doi:10.1007/BF02426650
33. van Westen GJ, Wegner JK, Geluykens P, Kwanten L, Vereycken I, Peeters A, IJzerman AP, van Vlijmen HW, Bender A (2011) Which compound to select in lead optimization? Prospectively validated proteochemometric models guide preclinical development. *PLoS One* 6 (11):e27518. doi:10.1371/journal.pone.0027518

Table S5.1: Atoms in vocabulary for graph-based molecule generation.

Symbol	Valence	Charge	Number	Word
O	2	0	8	2O
O+	3	1	8	3O+
O-	1	-1	8	1O-
C	4	0	6	4C
C+	3	1	6	3C+
C-	3	-1	6	3C-
N	3	0	7	3N
N+	4	1	7	4N+
N-	2	-1	7	2N-
Cl	1	0	17	1Cl
S	2	0	16	2S
S	6	0	16	6S
S	4	0	16	4S
S+	3	1	16	3S+
S+	5	1	16	5S+
S-	1	-1	16	1S-
F	1	0	9	1F
I	1	0	53	1I
I	5	0	53	5I
I+	2	1	53	2I+
Br	1	0	35	1Br
P	5	0	15	5P
P	3	0	15	3P
P+	4	1	15	4P+
Se	2	0	34	2Se
Se	6	0	34	6Se
Se	4	0	34	4Se
Se+	3	1	34	3Se+
Si	4	0	14	4Si
B	3	0	5	3B
B-	4	-1	5	4B-
As	5	0	33	5As
As	3	0	33	3As
As+	4	1	33	4As+
Te	2	0	52	2Te
Te	4	0	52	4Te
Te+	3	1	52	3Te+
*	0	0	0	*

The column of 'Symbol' is the symbol of the atom and its charge; the column of 'Valence' is the value of valence of the state of each chemical element; the 'Number' column stands for the index of each element in the periodic table, the last row is the unique word for each state of these elements, a combination of its valence and symbol.

Table S5.2: The pseudo code for encoding the graph representation of molecules in *DrugEx v3*

Algorithm encoding:

Input:

mol: structure of the kekulized molecule

subs: structure of the scaffolds

vocab: vocabulary of tokens which is consisted of graph matrix

Output:

matrix: the n x 5 matrix to represents the molecular graph.

```
# Ensure the atom of the subs are put at the start in the molecule
mol ← RANK_ATOM_BY_SUB(mol, subs)
sub_atoms ← GET_ATOMS (subs)
sub_bonds ← GET_BONDS (subs)
mol_atoms ← GET_ATOMS (mol)
frag, grow, link ← [('GO', 0, 0, 0, 1)], [], [(0, 0, 0, 0, 0)]
For atom in mol_atoms:
    # The bonds which connect to the atom having the index before this atom
    bonds ← GET_LEFT_BONDS (mol, atom)
    For bond in bonds:
        tk_bond ← GET_TOKEN (vocab, bond)
        other ← GET_OTHER_ATOM(mol, atom, bond)
        If IS_FIRST (bonds, bond):
            tk_atom ← GET_TOKEN (vocab, atom)
        Else:
            tk_atom ← GET_TOKEN (vocab, None)

        # The index of the scaffold in which the current atom locates
        # Its value starts from 1. If it is not in the scaffold, it will be 0
        scf ← GET_FRAG_ID (subs, atom)
        column ← (tk_atom, tk_bond, GET_INDEX (other), GET_INDEX (atom), scf)
        If other in sub_atoms and atom in sub_atoms and bond not in sub_bonds:
            Insert column to link
        Else if bond in sub_bonds:
            Insert column to frag
        Else:
            Insert column to grow
    End
End
Insert ('EOS', 0, 0, 0, 0) to grow
matrix ← CONCATENATE_BY_COLUMN (frag, grow, link)

Return matrix
```

Table S5.3: The pseudo code for decoding the graph representation of molecules in *DrugEx* v3

Algorithm decoding:

Input:

matrix: the $n \times 5$ matrix to represents the molecular graph

vocab: vocabulary of tokens which is consisted of graph matrix

Output:

mol: structure of the kekulized molecule

subs: structure of the scaffolds

```

mol ← new MOL ()
subs ← new SUB ()
For atom, bond, prev, curr, scf in matrix:
  If atom == 'EOS' or atom == 'GO':
    continue
  If atom != '*':
    a ← new Atom (GET_ATOM_SYMBOL(vocab, atom))
    SET_FORMAL_CHARGE (a, GET_CHARGE(vocab, atom))
    ADD_ATOM (mol, a)
    If scf != 0: ADD_ATOM (subs, a)
  If bond != 0:
    b ← new Bond (bond)
    ADD_BOND(mol, b)
    If frag != 0:
      ADD_BOND (subs, b)
End

# automatically determine the aromatic rings
mol ← SANITIZE (mol)
subs ← SANITIZE (subs)

Return mol, subs

```
