

RESEARCH

Open Access



DrugDiff: small molecule diffusion model with flexible guidance towards molecular properties

Marie Oestreich^{1,2*}, Erinc Merdivan³, Michael Lee^{1,2}, Joachim L. Schultze^{2,4,5}, Marie Piraud³ and Matthias Becker^{1,2*}

Abstract

With the cost/yield-ratio of drug development becoming increasingly unfavourable, recent work has explored machine learning to accelerate early stages of the development process. Given the current success of deep generative models across domains, we here investigated their application to the property-based proposal of new small molecules for drug development. Specifically, we trained a latent diffusion model—*DrugDiff*—paired with predictor guidance to generate novel compounds with a variety of desired molecular properties. The architecture was designed to be highly flexible and easily adaptable to future scenarios. Our experiments showed successful generation of unique, diverse and novel small molecules with targeted properties. The code is available at <https://github.com/MarieOestreich/DrugDiff>.

Scientific Contribution This work expands the use of generative modelling in the field of drug development from previously introduced models for proteins and RNA to the here presented application to small molecules. With small molecules making up the majority of drugs, but simultaneously being difficult to model due to their elaborate chemical rules, this work tackles a new level of difficulty in comparison to sequence-based molecule generation as is the case for proteins and RNA. Additionally, the demonstrated framework is highly flexible, allowing easy addition or removal of considered molecular properties without the need to retrain the model, making it highly adaptable to diverse research settings and it shows compelling performance for a wide variety of targeted molecular properties.

Keywords Drug development, Latent diffusion, Generative modelling, Targeted generation

*Correspondence:

Marie Oestreich
Marie.Oestreich@dzne.de
Matthias Becker
Matthias.Becker@dzne.de

¹ Modular High-Performance Computing and Artificial Intelligence,
German Center for Neurodegenerative Diseases (DZNE), Bonn, Germany

² Systems Medicine, German Center for Neurodegenerative Diseases
(DZNE), Bonn, Germany

³ Helmholtz AI, Helmholtz Munich, Ingolstädter Landstraße 1,
85764 Neuherberg, Germany

⁴ Genomics & Immunoregulation, Life and Medical Sciences (LIMES)
Institute, University of Bonn, Bonn, Germany

⁵ PRECISE Platform for Single Cell Genomics and Epigenomics, DZNE
and University of Bonn, Bonn, Germany



© The Author(s) 2025. **Open Access** This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

Introduction

The drug development process has become increasingly unsustainable over the last years, with the ratio of product yield to development costs becoming more and more unfavourable [1–4]. In an attempt to revert this trend, we can observe the development of generative machine learning models in recent years [5–13]. Particular advancements have been made recently with diffusion models for AI-based protein design [7, 8, 11]. However,

the drug space, i.e. the parts of the chemical space harvested for drug development, not only comprises proteins. The three dominant chemical subspaces are small molecules, proteins and nucleic acids such as RNA. With small molecules making up the majority of drugs [14], we developed *DrugDiff*, a diffusion model for the generation of small molecules (Fig. 1A), to help navigate the vast space of potentially drug-like small molecules. Several aspects suggested diffusion models as suitable candidates

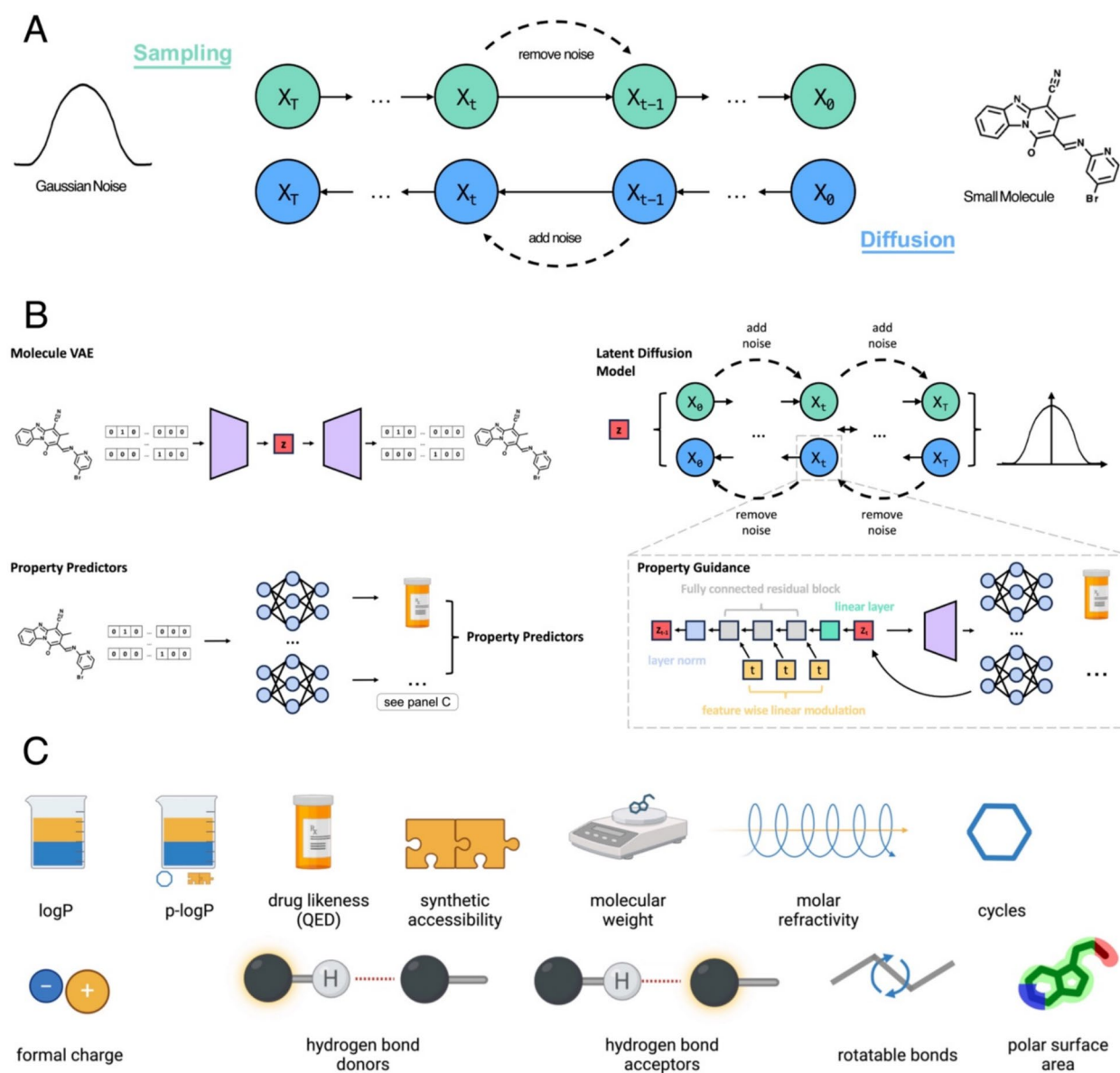


Fig. 1 DrugDiff Overview. **A** schematic of *DrugDiff*, which uses a diffusion model to generate molecules from Gaussian noise. Both the forward and the backward diffusion process are illustrated; **B** the detailed architecture of *DrugDiff*, which comprises a VAE (top left) that autoencodes one-hot-encoded SELFIES and whose latent space (z) serves as input to the latent diffusion model (top right). The bottom left illustrates molecular property predictors that are trained on one-hot encoded SELFIES and then used for guidance during the diffusion steps. The guidance together with the detailed architecture of the latent diffusion model is illustrated in the inset (bottom right); **C** The molecular properties used for guidance

for this task: (i) their recent successes not only in the image but also the protein domain demonstrate their potential to model highly complex data; (ii) when designing novel small molecule drugs, it is essential to offer guidance towards desired molecular properties during the generation process. To maximise utility and applicability in diverse scenarios, the set of properties to guide for must be highly flexible. With predictor-based guidance, diffusion models offer guided generation without explicit conditional training, introducing the required flexibility; (iii) unlike other generative model architectures, diffusion models do multi-step rather than one-shot generation. This allows gradual guidance towards desired properties with the option to correct missteps along the way and (iv) they can be combined with variational autoencoders (VAEs) to form a latent diffusion model [15]. Latent diffusion models are not trained on the original data space, but instead on a latent representation stemming from a pre-trained VAE. Latent diffusion models are particularly attractive in the context of small molecules, to address the often-raised topic of molecule representation, because unlike proteins and nucleic acids, small molecules are not chains of pre-existing building blocks like amino acids or nucleobases that are connected in a clearly defined manner. Instead, there are complex chemical rules that dictate, for instance, each atom's number of eligible bonds, possible charges or bond angles and these rules are difficult to represent. While several molecular representations exist [16], many of them are either discrete, not unique, or very sparse, whereas a continuous numeric representation is the preferred input for many deep learning models. Latent diffusion models outsource the task of learning a mapping from the molecular representation to a continuous latent space to a VAE, while the training of the diffusion model itself is focussed on modelling its latent distribution. Previous work has addressed diffusion models for small molecule generation, many of these focusing on diffusing three-dimensional (3D) atom coordinates [17–20]. 3D-diffusion methodologies for small molecules have demonstrated significant utility, particularly in the design of ligands tailored to fit into protein binding pockets [18]. However, this approach necessitates predefined specifications, such as the number of atoms, which imposes substantial constraints and requires considerable user input. Existing methods focus primarily on the target pocket [18], or they explicitly integrate specific molecular properties into the training process [20]. This design choice reduces flexibility, because it limits the adaptability of these models to properties that may not have been anticipated or explicitly encoded during development. A 3D-diffusion approach proposed by Luo et al. utilises text prompts to guide toward molecules with desired properties [19].

Though this introduces flexibility regarding available properties, it requires explicit descriptions of the molecular requirements. While some molecular properties can be effectively described using text prompts, others are more challenging to articulate, particularly spatially complex requirements such as fitting a molecule into a specific binding pocket. Moreover, this approach inherently biases generation toward pre-existing user preconceptions and expectations of molecular design, potentially limiting the process from producing unexpected or novel compounds. Other work has focused on the use of *latent* diffusion for small molecules, like we are proposing here. However, these approaches are primarily designed for the optimisation of existing molecules rather than *de novo* generation and do not address scenarios involving multi-property optimisation [21]. Others employ node-based latent representations, again necessitating the predefined specification of the number of nodes prior to generation [22]. This approach also involves direct conditioning during training, which constrains flexibility.

Our proposed model *DrugDiff* maintains flexibility with respect to target properties and minimises required user specification prior to generation as for example the number of atoms in the generated molecule. It comprises three parts (Fig. 1B): (1) A VAE trained on SELFIES representations of small molecules [23] from the public ZINC250K dataset, which contains approximately 250,000 small, drug-like and commercially available molecules from the ZINC database [24, 25]; (2) a latent diffusion model trained on the VAE's latent space and (3) a series of molecular property predictors trained independently on one-hot encodings of the ZINC250K molecules. The diffusion model learns to generate latent representations of small molecules by starting from pure Gaussian noise and predicting—in a series of diffusion steps—small amounts of noise to remove and gradually de-noise the latent representation. To guide this diffusion process towards desired molecular properties, the latents are decoded at every step, the decoded molecules are then passed to the pre-trained property predictors and the computed loss between desired and actual property is back-propagated directly onto the latent space. The latent space is then manipulated using the computed gradients before entering the next denoising step. We chose this method to guide towards molecular properties because we decided to explicitly avoid conditional training of the diffusion model. The consequence of conditional training would be that whenever new properties are to be added in future applications, the diffusion model would have to be retrained. Not only would that make the model inflexible in its direct application to future use cases, but it would also render the model very unsustainable. We addressed these issues by using predictor guidance. To illustrate the

successful generation of small molecules with defined molecular features with *DrugDiff*, we selected a variety of molecular properties that are relevant for the mechanism of action of drug-like small molecules (Fig. 1C).

Results

Unguided generation of small molecules

Before steering the generation process towards specific molecular properties, we first validated that the model had correctly learned from its training data distribution. To this end, we generated 10,000 molecules with *DrugDiff* without any guidance and additionally sampled the same number of molecules from the VAE's latent space, the distribution that our diffusion model was trained to learn. Supplementary Fig. 1A shows a random subset of the molecules generated with *DrugDiff*. The molecules exhibit diverse structures, featuring various elements, ring sizes, bond types and molecular sizes. Supplementary Fig. 1B quantifies *DrugDiff*'s ability to produce unique, novel and chemically diverse molecules. The scores are generally high and comparable to that of the VAE. The internal diversity score measures how chemically diverse a set of molecules is on a scale from 0 (not diverse) to 1 (very diverse), with a high score of 0.912 indicating *DrugDiff* did not only learn subspaces of the latent space but covered its full chemical information landscape. The high novelty and uniqueness scores show that *DrugDiff* was able to generalise to the underlying data distribution and neither suffers from mode collapse nor overfitting. We additionally evaluated the VAE and *DrugDiff* using the GuacaMol benchmarking dataset and associated Distribution-Learning Benchmarks [26]. To this end, we re-trained the VAE on the benchmarking set and subsequently re-trained *DrugDiff* on the new VAE latent space. The benchmark evaluation (Suppl. Table 1) shows that the results for VAE and *DrugDiff* are very similar, further underlining that *DrugDiff* is indeed capable of fully learning the VAE's latent space. We subsequently investigated in more detail whether *DrugDiff* could cover the full range of various molecular properties as they occurred in the training data. Supplementary Fig. 1C shows the distribution of 15 properties in the molecules generated with the VAE compared to those generated by *DrugDiff*. Those properties are: The topological polar surface area (a measurement for passive transport through membranes [27]), synthetic accessibility (a score to estimate how easily a molecule can be synthesised [28]), the number of rotatable bonds (single bonds that are not part of a ring and that are attached to an atom that is neither hydrogen nor terminal [29]), quantitative estimation of drug-likeness [30], molar refractivity (a measurement for how polarisable a molecule is [31]), molecular weight, the number of atoms and heavy atoms in particular, the logP

(indicates a molecule's lipophilicity [31]) as well as penalised logP (logP penalised for poor synthetic accessibility and large cycles [32]), the number of very small or large cycles (≤ 4 or ≥ 7 atoms), the number of hydrogen-bond donors and -acceptors, the number of cycles of any size, and formal charge. For all these properties, the property distributions of the molecules generated by *DrugDiff* and those randomly sampled from the latent space are highly overlapping.

In summary, these results illustrate that *DrugDiff* learned to cover the full distribution of the training data, and without guidance towards target properties it is capable of generating novel, unique and chemically diverse small molecules, demonstrating to good generalisation to the data space.

Single-property guidance

For a model to truly facilitate drug development by proposing novel small molecules, it *must* be able to generate such molecules under provided property constraints. While unconditional generation with subsequent filtering for molecular properties is theoretically an option, the lack in targeted generation typically results in the loss of most generated molecules after filtering. Such an approach would require a much larger number of molecules to be generated to account for the filter loss and molecules with property values residing at the tails of their distributions will be underrepresented. Actively guiding the generative process towards a desired property is therefore important to maximise the yield of molecules in the acceptable property range. Accordingly, we investigated *DrugDiff*'s ability to respond to guidance by a property predictor (for details, see methods). For each property, we first generated 10,000 molecules *without* guidance to use their property distribution as a point of reference. We then introduced guidance to the generation process for both directions, i.e. to increase and to decrease the property. To further illustrate the responsiveness of *DrugDiff* to the guidance signal, we used seven different guidance strengths when guiding towards an increased or decreased property value, generating 10,000 molecules in each case. Guidance strength refers to the factor by which the predictor loss is amplified and guidance strength of zero represents generation without guidance. Figure 2 illustrates the property distributions of the generated molecules at different guidance strengths and for different molecular properties. The properties that were investigated here are (from left to right and top to bottom): logP, synthetic accessibility, quantitative estimation of drug-likeness, penalised logP, molecular weight, molar refractivity, number of cycles (i.e. rings), number of hydrogen bond donors, formal charge, number of rotatable bonds, number of hydrogen bond acceptors

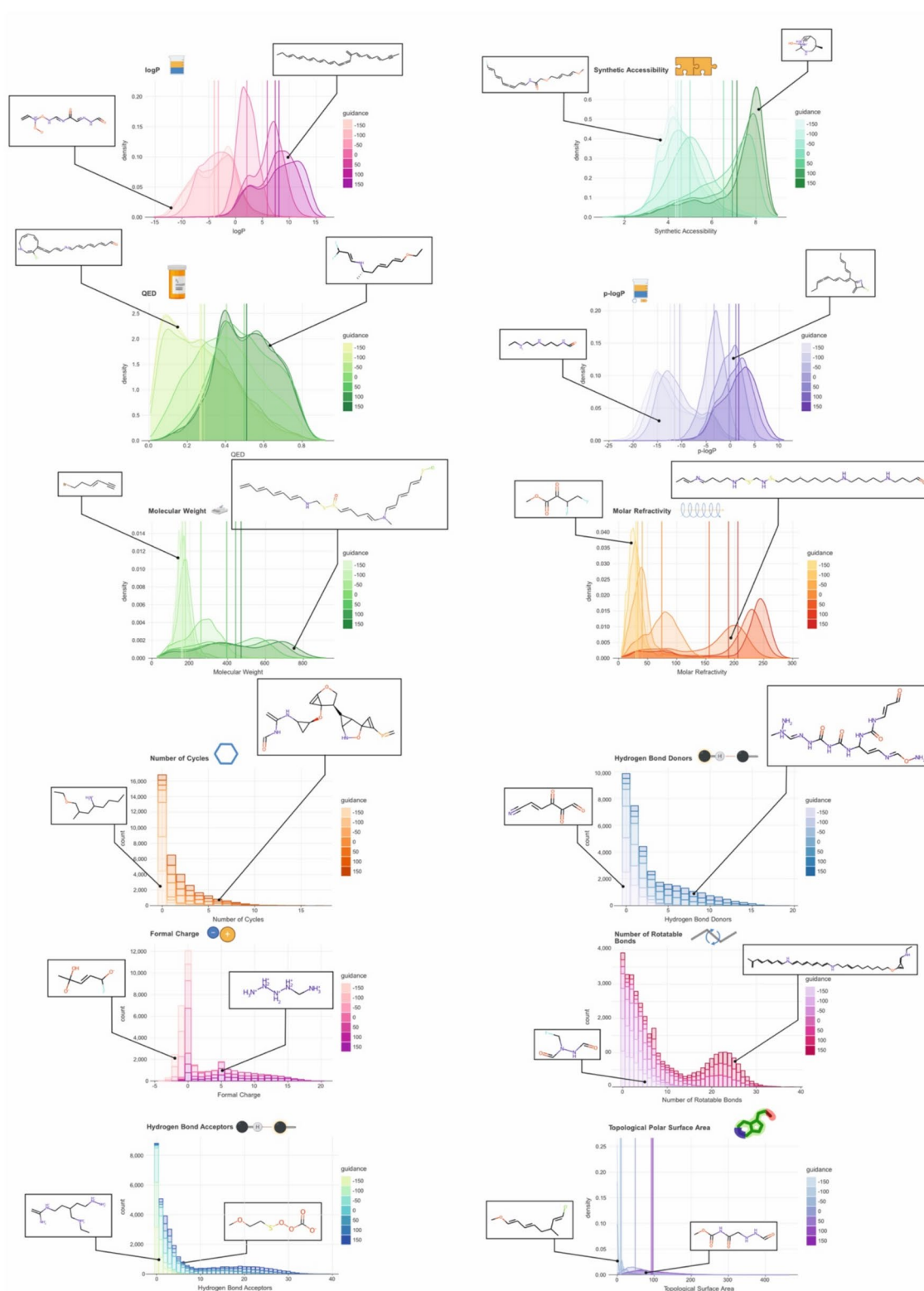


Fig. 2 Generated Ligands for ESR1. Depicted are the top 5 potential ligands generated with *DrugDiff* for the Human Estrogen Receptor (*ESR1*) based on their estimated free binding energy to the target. For each ligand, the SMILES and the free binding energy in kcal/mol estimated with AutoDock-GPU is given

In addition to the properties described above, another common goal during drug development is the generation of new ligands for a selected target protein. We thus additionally investigated *DrugDiff*'s ability to generate ligands with low estimated free binding energy using the Human Estrogen Receptor *ESR1* as a target. *ESR1* is a well-studied receptor known for its involvement in breast cancer

During drug development, it is often of particular interest to not only modify a single molecular property, but instead manipulate multiple properties in parallel. This task may include increasing or decreasing the different properties together or steering their values into opposing directions, reducing one while increasing another. We therefore next investigated *DrugDiff*'s ability to accommodate signals of several property predictors during the generation process for the purpose of multi-property guidance. Specifically, we considered the logP as well as the number of heavy atoms and therefore the overall size of the molecules. We considered four scenarios: (1) low logP and low number of atoms, (2) high logP and low number of atoms, (3) low logP and high number of atoms and (4) high logP and high number of atoms. For each scenario we generated 10,000 molecules. As can be seen in Fig. 4A, the density of molecules generated in scenario 1 (teal) is indeed highest in regions of low logP and low

Fig. 3 Single-Property Guidance. Shown are the distributions of various molecular properties in a set of 10,000 generated molecules when guiding the generation process with different guidance strengths. A guidance strength of zero equates to unguided generation, positive guidance strengths intend to shift the distribution towards higher values and negative guidance strengths towards lower values. Vertical lines represent the mean property value for each guidance strength. For illustration purposes, a generated molecule has been selected at random for strong positive and strong negative guidance and plotted alongside the distribution

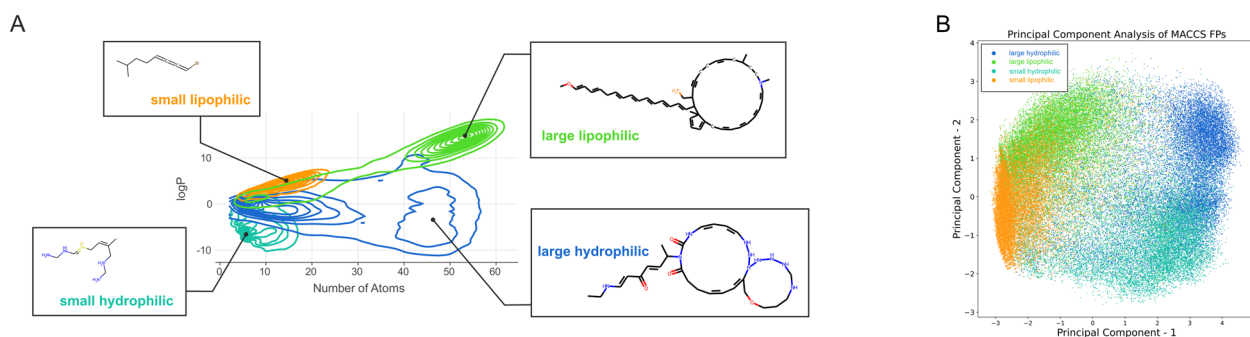


Fig. 4 Multi-Property Guidance. **A** Shown are the distributions of molecules in terms of logP and the number of heavy atoms generated with dual guidance in four different scenarios: high logP and low number of atoms (orange), low logP and low number of atoms (teal), high logP and high number of atoms (green), low logP and high number of atoms (blue). For each distribution, a molecule was picked at random and illustrated alongside it; **B** Shown are the first two principal components of a PCA conducted on MACCS fingerprints of the molecules from the four aforementioned scenarios. The colour-coding corresponds to that in **A**

number of atoms. Similarly, molecules generated in scenario 2 (orange) show much higher, positive logP values than those in scenario 1, while their number of atoms remains low. While in scenario 4 (green), a small subset is mislocated at low number of atoms the majority of molecules indeed have both high logP and number of atoms. Finally, scenario 3 (blue)—low logP and high number of atoms—was the hardest for the model to realise. While there is a population that fulfils these conditions, the large majority is located at lower heavy atom counts. Figure 4B shows a principal component analysis (PCA) with the first two principal components (PCs) displayed for the molecules generated in the different scenarios. The PCA was computed on the molecules' MACCS fingerprints (rdkit v2022.09.5), which indicate the presence of predefined substructures in a molecule. Clear spatial separation can be observed between the molecules generated for the different scenarios, with PC1 separating lipophilic from hydrophilic molecules and PC2 separating small from large.

Discussion

In this work, we have introduced a latent diffusion model with property guidance for the generation of novel small molecules with desired target properties. Existing generative models for small molecules are rigid in their design and do not allow easy adaptation to different use cases. This particularly concerns the way molecular properties are included in the model design. While the majority of works include them as conditions during the training process, we decided against this approach since it requires retraining of the entire model whenever properties are to be exchanged or added. Instead, *DrugDiff* utilises predictor guidance, where only chemical property predictors need to be trained and plugged in to include additional properties, but the generative model—which

is much more expensive in terms of training resources than a predictor – does not require retraining. Further, by using a latent diffusion model, the VAE can be easily replaced in the framework for one that handles a different modality than SELFIES, while the overall setup of the diffusion model and its training procedure remain unchanged. The conducted experiments clearly demonstrate that *DrugDiff* (i) was able to learn the breadth of its training distribution, (ii) can distinctly manipulate the molecular properties of molecules during the generation process, including complex properties such as the free binding energy to a target protein, and (iii) is capable to also expand this to a multi-property setting. In the context of multi-property guidance, it is important to highlight that the simultaneously applied properties must be reasonable. As could be seen on the example of decreasing logP—i.e. increasing hydrophilicity—and at the same time increasing the number of heavy atoms—i.e. making the molecules larger—it was difficult to produce a large population of molecules that fulfilled both criteria. Given that large organic substances are predominantly made up of a lipophilic hydro-carbon backbone, difficulties in creating large organic molecules with very low logP are to be expected. Hence, when using multi-property guidance, the desired conditions must be chosen to fit realistic scenarios in order to ensure a high yield.

When comparing property optimisation to other SOTA models, *DrugDiff*'s performance was comparable. The reason why the synthetic accessibility optimisation was more successful than the drug likeness is suspected to be the quality of the property predictor used: While the SA predictor had an r-value of 0.97 between target- and prediction-values on a hold-out test set, the QED predictor achieved an r-value of 0.83. Hyperparameter optimisation of this predictor may improve its prediction accuracy and would be expected to improve the optimisation

success. Given that the predictors serve as experts to the diffusion model which it consults for feedback during generation, predictor performance is expected to impact property optimisation performance.

In summary, *DrugDiff* is a diffusion model for the generation of small molecules with desired target properties, which was built in a highly flexible framework. The modular architecture allows high customisability, allowing the exchange of the VAE including the molecular representations it receives, as well as the properties used for guidance, making it easily adaptable to diverse application scenarios while minimising the required retraining of the model. Future expansions should include more complex properties, such as binding selectivity or systemic biological effects. Another future direction is training on much larger datasets to cover a vaster chemical space. While good results could already be achieved by training on the relatively small training set ZINC250K, more training data can be expected to further push the performance. Since the diffusion model learns from the VAE's latent space, latent space quality is instrumental for the diffusion model's success. Optimising the VAE with respect to the representation of chemistry in the latent space can therefore be expected to further improve performance. Additionally, a common task in drug development pipelines is substructure optimisation. Here, a substructure is provided as a starting point and modified to improve selected molecular properties. Incorporating this into the *DrugDiff* framework will add another layer of utility. Lastly, an end-to-end experiment starting with in silico generation of leads with *DrugDiff*, followed by synthesis and in vitro and/or in vivo evaluation will give insight into the exact speed-up of the development process offered by the model.

Methods

The model

The model constitutes three parts, a variational autoencoder that maps SELFIES onto a latent representation, the diffusion model that is trained to generate such molecule latents and a collection of property-predictors that are used to guide the diffusion process.

VAE

Variational Autoencoders [35] comprise two models, an encoder and a decoder. The encoder $g(x) = z$ is used to embed higher-dimensional data onto a continuous-valued latent space of typically smaller size. The decoder $f(z) = x' \approx x$ is trained to revert the embedding back to the original input. The VAE used for the embedding of molecules for further processing by the diffusion model was taken from Eckmann et al. [23].

The model was pre-trained on one-hot encodings of SELFIES from the ZINC250K dataset, a subset of the ZINC database [24, 25, 36], comprising approximately 250,000 small molecules that fulfil the Lipinski Rule of 5 [37]. The model utilises a latent size of 1024. Training the VAE was done following the steps described in the repository accompanying the paper [23].

Diffusion model

Our diffusion model is a *latent diffusion model* based on the implementation by Rombach et al. [15]. More precisely, the underlying model architecture is a *denoising diffusion probabilistic model* (DDPM) [38] with a linear layer, three fully connected residual blocks, a layer norm and another linear layer. Timesteps are embedded and provided to the fully connected residual blocks via feature-wise linear modulation layers.

DDPMs are trained to revert a Markov chain of noising steps, the diffusion process. This diffusion process incrementally adds Gaussian noise to a sample over $0, \dots, T$ timesteps such that a sample x_0 from the original domain is transitioned into x_T , which is Gaussian noise for large T . The noising step at t is conditioned on the noised sample x_{t-1} of the previous timestep [38]

$$q(x_t|x_{t-1}) = N\left(x_t; \sqrt{1 - \beta_t}x_{t-1}, \beta_t I\right)$$

where t is taken from a pre-defined, time-dependent noising schedule. In this formulation, the computation of x_t first requires the computation of all previous timesteps $x_{0:t-1}$, however, applying the reparameterization trick [35, 38], it can be directly conditioned on x_0 with

$$q(x_t|x_0) = N\left(x_t; \sqrt{\alpha_t}x_0, \sqrt{1 - \alpha_t}I\right)$$

where $\alpha_t = 1 - \beta_t$ and $\bar{\alpha}_t = \prod_{i=1}^t \alpha_i$. The noising schedule used here is linear. The DDPM is trained to reverse this diffusion process and therefore generate a sample from $q(x_{t-1}|x_t)$. Stepwise application for $t = T, \dots, 0$ then allows the generation of samples from the original domain, $x_0 \sim q(x_0)$, starting from Gaussian noise x_T . In this work, we are rephrasing the sampling step as a noise predictor that predicts t rather than x_{t-1} directly, as proposed by Ho et al. [38], which is trained using

$$L_t = \mathbb{E}_{t, x_0, \epsilon} \left[\|\epsilon - \epsilon_\theta\left(\sqrt{\alpha_t}x_0 + \sqrt{1 - \bar{\alpha}_t}\epsilon, t\right)\|^2 \right]$$

for $0 < t \leq T$. Note that $\sqrt{\alpha_t}x_0 + \sqrt{1 - \bar{\alpha}_t}\epsilon = x_t$ after the reparameterization trick. The latent diffusion model was trained for 100 epochs on the VAE-embeddings of the ZINC250K data.

Property predictors

We use a series of property predictors to guide the diffusion model during sampling. The property predictors followed the implementation of Eckmann et al. [23] and were trained on the ZINC250K dataset. Drug-likeness (QED), logP, molecular weight, molar refractivity, topological polar surface area, number of H-bond acceptors/-donors, number of rotatable bonds, number of atoms, formal charge and number of rings were all computed using *rdkit* (v2022.09.5). The synthetic accessibility was calculated using the *sascorer.py* from [23] which is based on [28]. The p-logP was computed as follows:

$$p - \log P = \log P - SA - n_{\text{larger rings}}$$

where $n_{\text{larger rings}}$ is the number of rings with more than 6 atoms. The free binding energies to the protein target *ESR1* were estimated using AutoDock-GPU version: v1.6-release [34]. With 5 docking runs per molecule and using the energy of the best binding pose as the target value. The maps used for docking were in accordance with Eckmann et al. [23]. All predictors were trained on decoded one-hot encodings of the molecules. During guidance, the noisy latents were first decoded by the VAE's decoder and then passed to the property predictors. We then computed the gradients based on the predicted properties and used them to manipulate the predicted noise ϵ_t .

Evaluation metrics

To evaluate the molecules generated by our latent diffusion model, we used three metrics: (1) novelty, the fraction of generated molecules that are not found in the training set, to ensure that the model does not simply learn to copy the molecules; (2) uniqueness, the percentage of unique molecules among all generated ones; (3) internal diversity, a metric implemented in the *molsets* library that assesses how chemically diverse the generated molecules are. With this metric we can detect mode collapse, i.e. the model producing only a set of highly similar molecules to satisfy a given property.

Supplementary Information

The online version contains supplementary material available at <https://doi.org/10.1186/s13321-025-00965-x>.

Supplementary material 1.
Supplementary material 2.
Supplementary material 3.

Acknowledgements

This work was supported by the HGF Helmholtz AI grant Pro-Gene-Gen ZT-I-PF5-23 as well as the Helmholtz Association's Initiative and Networking Fund on the HAICORE@FZJ partition. Further support came from the EU supported

BMBF grant PriSyn (16KISA030). Figures were partially created using Biorender.com.

Author contributions

MO, EM, MP and MB conceptualized the study, MO, EM and ML conducted the experiments, MO, EM, ML, JLS, MP and MB wrote the manuscript. All authors read the manuscript, provided feedback and eventually approved it in its final form.

Funding

Open Access funding enabled and organized by Projekt DEAL. This work was supported by the HGF Helmholtz AI grant Pro-Gene-Gen ZT-I-PF5-23 as well as the Helmholtz Association's Initiative and Networking Fund on the HAICORE@FZJ partition. Further support came from the EU supported BMBF grant PriSyn (16KISA030). Figures were partially created using Biorender.com.

Availability of data and materials

All data used stems from publicly available databases.

Code availability

The source code, model checkpoints and notebooks for reproducing the results presented here can be found on GitHub (<https://github.com/MarieOestreich/DrugDiff>) and Zenodo (<https://doi.org/10.5281/zenodo.12755762>).

Declarations

Ethics approval and consent to participate

An issue that needs addressing is the dual-use problem of models like the one presented here. Generating molecules that manifest desired properties has many beneficial use cases, for example in the medical context or materials sciences. However, it also bears the inherent risk of being misused for malicious intent, for example to generate molecules that maximise toxicity. The model described in this work is solely intended to be used in a beneficial context and in accordance with the law.

Competing interests

The authors declare no competing interests.

Received: 24 October 2024 Accepted: 27 January 2025

Published online: 25 February 2025

References

- Sadybekov AV, Katritch V (2023) Computational approaches streamlining drug discovery. *Nature* 616(7958):673–685. <https://doi.org/10.1038/s41586-023-05905-z>
- Deloitte pharma study: Drop-off in returns on R&D investments—sharp decline in peak sales per asset. <https://www2.deloitte.com/ch/en/pages/press-releases/articles/deloitte-pharma-study-drop-off-in-returns-on-r-and-d-investments-sharp-decline-in-peak-sales-per-asset.html> (accessed May 08, 2023).
- Deloitte. (2023). Seize the digital momentum Measuring the return from pharmaceutical innovation 2022.
- Scannell JW, Blanckley A, Boldon H, Warrington B (2012) Diagnosing the decline in pharmaceutical R&D efficiency. *Nat Rev Drug Discov* 11(3):191–200. <https://doi.org/10.1038/nrd3681>
- Ren F et al (2023) AlphaFold accelerates artificial intelligence powered drug discovery: efficient discovery of a novel CDK20 small molecule inhibitor. *Chem Sci* 14(6):1443–1452. <https://doi.org/10.1039/d2sc05709c>
- Zhavoronkov A et al (2019) Deep learning enables rapid identification of potent DDR1 kinase inhibitors. *Nat Biotechnol* 37(9):1038–1040. <https://doi.org/10.1038/s41587-019-0224-x>
- Watson JL et al (2023) De novo design of protein structure and function with RFdiffusion. *Nature* 620(7976):1089–1100. <https://doi.org/10.1038/s41586-023-06415-8>

8. Abramson J et al (2024) Accurate structure prediction of biomolecular interactions with AlphaFold 3. *Nature* 630(8016):493–500. <https://doi.org/10.1038/s41586-024-07487-w>
9. Guo J et al (2023) Link-INVET: generative linker design with reinforcement learning. *Digital Discovery* 2(2):392–408. <https://doi.org/10.1039/D2DD000115B>
10. Imrie F, Hadfield TE, Bradley AR, Deane CM (2021) Deep generative design with 3D pharmacophoric constraints. *Chem Sci* 12(43):14577–14589. <https://doi.org/10.1039/d1sc02436a>
11. Ingraham J et al (2022) Illuminating protein space with a programmable generative model. *BioRxiv*. <https://doi.org/10.1101/2022.12.01.518682>
12. Schaller D et al (2020) Next generation 3D pharmacophore modeling. *WIREs Comput Mol Sci*. <https://doi.org/10.1002/wcms.1468>
13. Torge J, Harris C, Mathis SV, Lio P (2023) DiffHopp: a graph diffusion model for novel drug design via scaffold hopping. *arXiv*. <https://doi.org/10.48550/arxiv.2308.07416>
14. Howes L (2023) Is this a golden age of small-molecule drug discovery? *C&EN Global Enterp* 101(36):28–32. <https://doi.org/10.1021/cen-10136-cover>
15. R. Rombach, A. Blattmann, D. Lorenz, P. Esser, and B. Ommer. (2021). High-resolution image synthesis with latent diffusion models. *arXiv*. <https://doi.org/10.48550/arxiv.2112.10752>
16. Du Y et al (2024) Machine learning-aided generative molecular design. *Nat Mach Intell*. <https://doi.org/10.1038/s42256-024-00843-5>
17. Alakhdar A, Poczos B, Washburn N (2024) Diffusion models in de novo drug design. *J Chem Inf Model* 64(19):7238–7256. <https://doi.org/10.1021/acs.jcim.4c01107>
18. Huang L et al (2024) A dual diffusion model enables 3D molecule generation and lead optimization based on target pockets. *Nat Commun* 15(1):2657. <https://doi.org/10.1038/s41467-024-46569-1>
19. Luo Y et al (2024) Text-guided small molecule generation via diffusion model. *iScience* 27:110992. <https://doi.org/10.1016/j.isci.2024.110992>
20. Qiang B et al (2023) Coarse-to-Fine: a hierarchical diffusion model for molecule generation in 3D. *Int Conf Machine Learn*. <https://doi.org/10.48550/arxiv.2305.13266>
21. Runcie NT, Mey ASJS (2023) SILVR: guided diffusion for molecule generation. *J Chem Inf Model* 63(19):5996–6005. <https://doi.org/10.1021/acs.jcim.3c00667>
22. Xu M, Powers A, Dror R, Ermon S, Leskovec J (2023) Geometric latent diffusion models for 3D molecule generation. *Int Conf Machine Learn*. <https://doi.org/10.48550/arxiv.2305.01140>
23. Eckmann P, Sun K, Zhao B, Feng M, Gilson MK, Yu R (2022) LIMO: Latent inceptionism for targeted molecule generation presented at the ICML. *Proc Machine Learn Res* 162:5777
24. Sterling T, Irwin JJ (2015) ZINC 15—ligand discovery for everyone. *J Chem Inf Model* 55(11):2324–2337. <https://doi.org/10.1021/acs.jcim.5b00559>
25. Irwin JJ et al (2020) ZINC20-A free ultralarge-scale chemical database for ligand discovery. *J Chem Inf Model* 60(12):6065–6073. <https://doi.org/10.1021/acs.jcim.0c00675>
26. Brown N, Fiscato M, Segler MHS, Vaucher AC (2019) Guacamol: benchmarking models for de novo molecular design. *J Chem Inf Model* 59(3):1096–1108. <https://doi.org/10.1021/acs.jcim.8b00839>
27. Ertl P, Rohde B, Selzer P (2000) Fast calculation of molecular polar surface area as a sum of fragment-based contributions and its application to the prediction of drug transport properties. *J Med Chem* 43(20):3714–3717. <https://doi.org/10.1021/jm000942e>
28. Ertl P, Schuffenhauer A (2009) Estimation of synthetic accessibility score of drug-like molecules based on molecular complexity and fragment contributions. *J Cheminform*. <https://doi.org/10.1186/1758-2946-1-8>
29. Veber DF, Johnson SR, Cheng H-Y, Smith BR, Ward KW, Kopple KD (2002) Molecular properties that influence the oral bioavailability of drug candidates. *J Med Chem* 45(12):2615–2623. <https://doi.org/10.1021/jm020017n>
30. Bickerton GR, Paolini GV, Besnard J, Muresan S, Hopkins AL (2012) Quantifying the chemical beauty of drugs. *Nat Chem* 4(2):90–98. <https://doi.org/10.1038/nchem.1243>
31. Wildman SA, Crippen GM (1999) Prediction of physicochemical parameters by atomic contributions. *J Chem Inf Comput Sci* 39(5):868–873. <https://doi.org/10.1021/ci990307i>
32. M. J. Kusner, B. Paige, and J. M. Hernández-Lobato. (2017). Grammar Variational Autoencoder.pdf," Proceedings of the 34th International Conference on Machine Learning.
33. Mazuz E, Shtar G, Shapira B, Rokach L (2023) Molecule generation using transformers and policy gradient reinforcement learning. *Sci Rep* 13(1):8799. <https://doi.org/10.1038/s41598-023-35648-w>
34. Santos-Martins D, Solis-Vasquez L, Tillack AF, Sanner MF, Koch A, Forli S (2021) Accelerating AutoDock4 with GPUs and gradient-based local search. *J Chem Theory Comput* 17(2):1060–1073. <https://doi.org/10.1021/acs.jctc.0c01006>
35. D. P. Kingma and M. (2013). Welling. Auto-encoding variational bayes. *arXiv*
36. Irwin JJ, Sterling T, Mysinger MM, Bolstad ES, Coleman RG (2012) ZINC: a free tool to discover chemistry for biology. *J Chem Inf Model* 52(7):1757–1768. <https://doi.org/10.1021/ci3001277>
37. Lipinski CA, Lombardo F, Dominy BW, Feeney PJ (2001) Experimental and computational approaches to estimate solubility and permeability in drug discovery and development settings. *Adv Drug Deliv Rev* 46(1–3):3–26. [https://doi.org/10.1016/S0169-409X\(96\)00423-1](https://doi.org/10.1016/S0169-409X(96)00423-1)
38. Ho J (2020) Denoising diffusion probabilistic models. *Adv Neural Inform Process Syst* 33:6480

Publisher's Note

Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.