

AI-enhanced adaptive virtual screening of large libraries for ligand discovery

Received: 8 May 2025

Accepted: 10 June 2026

Published online: 01 September 2026

 Check for updates

A list of authors and their affiliations appears at the end of the paper

Ultralarge virtual screenings (ULVSs) evaluate billions of molecules for drug discovery but face cost, flexibility and scalability limits. We introduce AdaptiveFlow, an open-source platform that makes ULVSs more accessible, scalable and efficient and supports artificial intelligence (AI) and machine learning (ML) method development. AdaptiveFlow provides a screening-ready version of the Enamine REAL Space, to our knowledge the largest library of ready-to-dock, drug-like molecules, comprising 69 billion compounds, also available in SELFIES format. An 18-dimensional grid of molecular properties prioritizes promising chemical subspaces, with optional active learning, reducing computational costs by orders of magnitude. AdaptiveFlow integrates >1,500 docking protocols, including GPU-accelerated and ML-based methods, and achieves near-linear scaling on up to 5.6 million CPUs in the Amazon Web Services cloud. We identified nanomolar inhibitors of two disease-relevant targets, ferroptosis suppressor protein 1 (FSP1) and poly(ADP-ribose) polymerase 1. Co-crystal structures provided mechanistic insights into FSP1 inhibition. AdaptiveFlow enables drug discovery at unprecedented scale and supports the development of AI-driven methods.

A central challenge in early-stage drug discovery is the identification and optimization of initial hit and lead compounds that bind specifically and potently to biological macromolecules. Historically, experimental high-throughput screening (HTS) has served as the primary method for discovering such hits¹. While these approaches have yielded many successful leads, they are inherently constrained by high costs, long timelines and the need for robust, scalable assays. These limitations are further compounded by the vastness of chemical space: the number of synthetically accessible, drug-like small molecules is estimated to exceed 10^{60} (ref. 2), yet typical HTS campaigns screen only hundreds of thousands of compounds. This minuscule sampling poses a major bottleneck, particularly for challenging targets such as protein–protein interaction interfaces or allosteric sites, where binding pockets are often shallow, dynamic or poorly defined. Virtual screenings have emerged as a compelling alternative, offering the ability to computationally evaluate extremely large numbers of molecules. Virtual screening surpassed millions of compounds over a decade ago³ and, in recent years, has expanded to hundreds of millions to billions of molecules in

silico, giving rise to ultralarge virtual screenings (ULVSs). Compared to traditional HTS, ULVSs dramatically reduce cost and time, while expanding the scope of chemical space that can be explored. In addition to identifying potent candidate compounds, virtual screening can provide insights into binding mechanisms and structure–activity relationships, making it especially valuable for targets that are difficult to modulate using conventional experimental approaches.

In recent years, multiple studies have demonstrated the remarkable effectiveness of ULVSs in identifying potent hits across a wide range of protein targets, including enzymatic active sites, orthosteric sites of G-protein-coupled receptors and protein–protein interaction interfaces^{4–8}. Driven by the success of billion-compound screens and advancements in synthesis and computational technologies, the number of available on-demand molecules with an established synthetic route has expanded rapidly, with current estimates reaching up to 3 trillion compounds across multiple vendors⁹.

While ULVSs enable the exploration of vast chemical libraries, notable challenges remain in making these searches efficient,

✉ e-mail: andrea.mattevi@unipv.it; hari@hms.harvard.edu; christoph.gorgulla@stjude.org

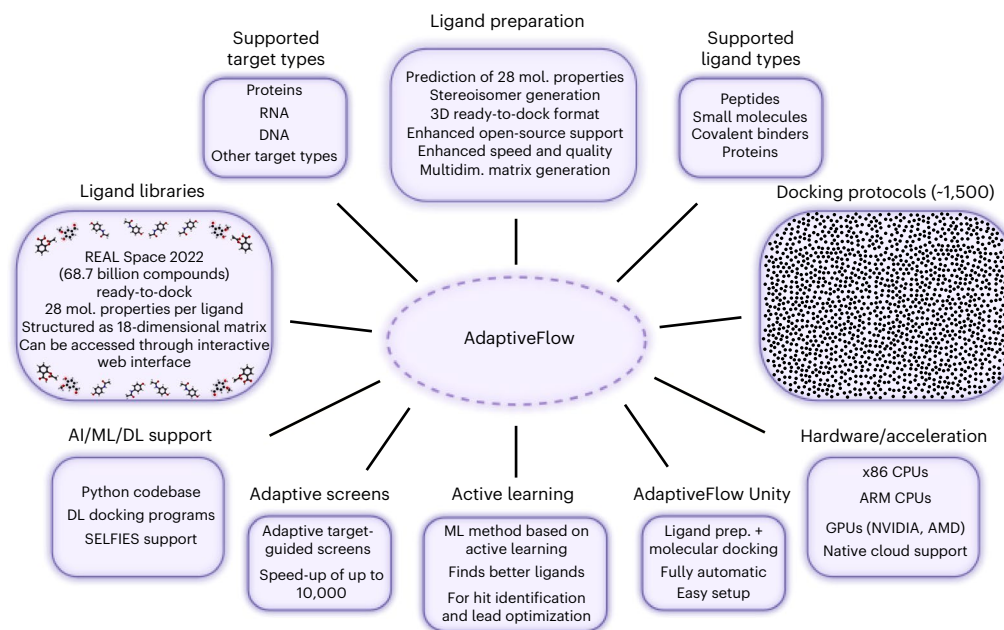


Fig. 1 | The AdaptiveFlow platform for ULVSs. AdaptiveFlow introduces several key advances that greatly expand the capabilities of ULVS, including (1) curation of the Enamine REAL Space library, comprising 69 billion molecules in a ready-to-dock format compatible in multiple chemical file formats; (2) integration of an extensive suite of approximately 1,500 docking protocols (defined as specific combinations of sampling algorithms and scoring functions), enabling the targeting of a broad range of receptor types, including RNA and DNA; (3) support for emerging hardware architectures such as ARM CPUs, along with GPU acceleration to enhance computational performance; (4) advanced ligand preparation (prep.) features, including stereoisomer

enumeration and the calculation of 28 molecular (mol.) properties, to enable the creation of more comprehensive screening libraries; and (5) built-in support for ML and deep learning approaches, including compatibility with SELFIES and neural-network-based docking protocols. The integrated module, AFU, further streamlines workflows by combining ligand preparation and docking into a single environment. Importantly, AdaptiveFlow also introduces a novel screening paradigm known as ATG-VSs, which enables efficient, target-focused exploration of ultralarge compound libraries at a fraction of the computational cost typically associated with conventional ULVSs. DL, deep learning; Multidim., multidimensional.

affordable and intelligent. Navigating such immense chemical space to rapidly identify promising hits requires additional strategies that go beyond brute-force docking. The financial cost alone is a major bottleneck⁸, rendering it impractical as chemical databases continue to grow exponentially, with newer libraries exceeding 1 trillion compounds and growing. Moreover, effective virtual screening requires the ability to flexibly integrate a wide range of docking algorithms, each built on different theoretical foundations and requiring distinct input formats, to better match the characteristics of diverse target classes. Compounding these challenges is the need to fully leverage modern computational infrastructure; while many existing platforms are limited to CPUs, there is increasing demand to exploit both CPU clusters and GPU architectures to scale performance and reduce turnaround time. These limitations—high cost, limited flexibility and insufficient scalability—highlight the urgent need for a new solution. To end this, we developed AdaptiveFlow, an open-source platform designed to intelligently navigate ultralarge chemical spaces, seamlessly incorporate diverse docking protocols and operate efficiently across heterogeneous computing environments, including cloud-scale CPU and GPU resources.

A key innovation of the platform is the implementation of Adaptive Target-Guided Virtual Screens (ATG-VSs), which use an 18-dimensional grid of molecular properties to identify and hone in on the most promising regions of chemical space, favorable for a chosen target site, substantially reducing computational cost relative to exhaustive docking. An optional active learning layer can be added to further refine and accelerate this process by iteratively selecting the most promising compounds of each tranche. To demonstrate the utility of AdaptiveFlow, we applied it to two therapeutically relevant targets: ferroptosis suppressor protein 1 (FSP1), a key regulator of lipid peroxidation and ferroptotic cell death with emerging roles in cancer resistance¹⁰, and poly(ADP-ribose) polymerase 1 (PARP1), a well-established DNA

repair enzyme and clinical target in oncology¹¹. In both cases, the platform identified nanomolar inhibitors, which were experimentally validated through high-resolution co-crystal structures confirming target engagement.

AdaptiveFlow

Deploying ULVSs at the scale of billions of compounds requires not only massive computational resources but also software systems that can efficiently coordinate the preparation, management and docking of libraries across heterogeneous high-performance computing environments. To meet these requirements, we developed AdaptiveFlow, an open-source platform built to streamline, scale and intelligently guide ULVS workflows from end to end. AdaptiveFlow combines deep flexibility with performance, offering support for diverse ligand and receptor types, GPU acceleration and modular integration of classical and machine learning (ML)-based docking protocols.

AdaptiveFlow introduces major enhancements across three core components: the AdaptiveFlow Ligand Preparation (AFLP) module for efficient preprocessing of massive libraries; the AdaptiveFlow for Virtual Screening (AFVS) engine, which supports over 1,500 docking protocols; and the AdaptiveFlow Unity (AFU) module, which integrates these components into a unified, modular workflow adaptable to both CPU-based and GPU-based infrastructures (Fig. 1 and Supplementary Table 1). These modules have been engineered to support both traditional and emerging use cases while optimizing for large-scale parallel execution on cloud and on-premise infrastructures. AdaptiveFlow is natively compatible with cloud services such as Amazon Web Services (AWS) (Supplementary Fig. 1), where it demonstrates near-linear scaling up to 5.6 million virtual CPUs (vCPUs) (Supplementary Fig. 2), enabling efficient exploration of ultralarge libraries. The use of spot (preemptible) instances further reduces the costs.

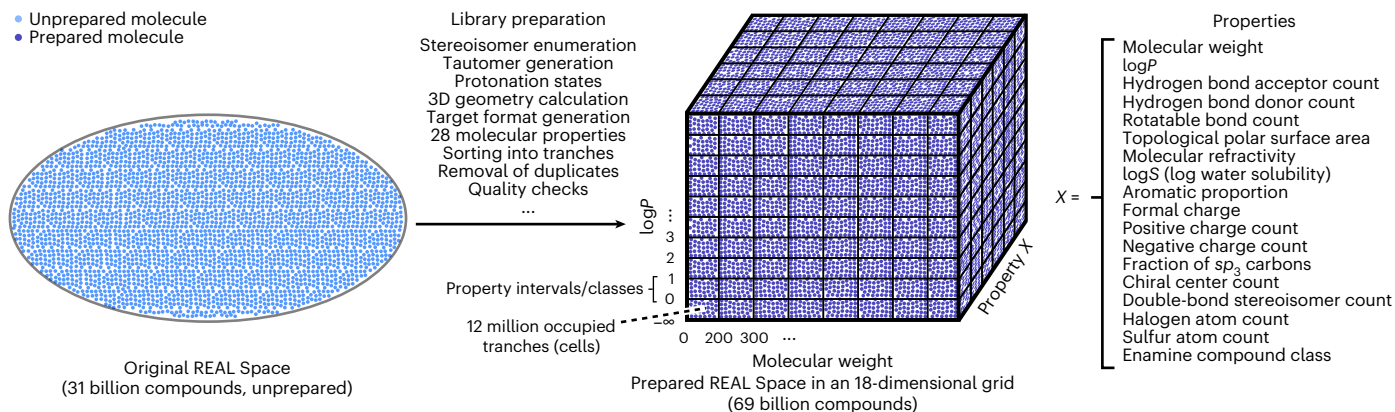


Fig. 2 | Organization and preparation of the Enamine REAL Space. The initial version of the REAL Space (enumerated release, 2022q1-2) comprised approximately 31 billion molecules, represented as SMILES strings in an unprocessed format (for example, without full stereoisomer enumeration or structural standardization). A selected subset of 18 of these properties were used to partition the library into an 18-dimensional property matrix, where each dimension was discretized into multiple intervals. In this figure, the

high-dimensional matrix is illustrated as a 3D schematic for clarity, with three representative properties each split into 6–10 intervals, resulting in $8 \times 6 \times 6 = 288$ tranches. The actual prepared REAL Space matrix has around 12 million occupied tranches. Each box represents a tranche that contains ligands sharing the corresponding property ranges. This structure enables flexible selection of compound subsets and serves as the foundation for ATG-VSs that are introduced in this work.

The AFLP module provides a pipeline for preparing massive compound libraries for docking. In addition to generating three-dimensional (3D) conformers, it enumerates stereoisomers and tautomers, validates geometries and calculates 28 molecular properties (Fig. 2, Extended Data Figs. 1–3, Supplementary Tables 2 and 3 and Supplementary Figs. 3–10). These properties are used to disperse the molecules within a multidimensional grid, up to 18 dimensions, on the basis of physicochemical descriptors. This grid-based representation forms the foundation for ATG-VSs, which allow users to intelligently select and prioritize chemically diverse subspaces relevant to a given target, thereby optimizing computational resources focusing on the most promising regions of chemical space. Furthermore, an optional active learning component can be layered on top of ATG-VSs to iteratively refine the search using data-driven feedback, dynamically selecting molecules predicted to yield the highest information gain or binding likelihood (Supplementary Figs. 1–3 and 11 and Extended Data Fig. 4). Because each screening stage generates new labeled data, this cycle can be repeated for an arbitrary number of rounds, retraining the model on the newly docked and, where available, experimentally assayed compounds and re-selecting the next batch to screen, thereby forming a closed-loop active-learning cycle that progressively concentrates sampling on the most promising regions of chemical space.

The AFVS module enables large-scale docking campaigns using over 1,500 distinct docking protocols. We define a docking protocol as a specific combination of a sampling algorithm (pose prediction) and a scoring function, allowing for a vast number of unique workflows assembled from more than 40 individual docking programs. AFVS supports a wide range of ligand and target types, including peptides, DNA and RNA, and includes recent advances in ML-based docking such as DiffDock and TANKBind^{12,13}, which have shown promise in speed and accuracy, especially for blind or flexible docking scenarios (Supplementary Tables 4–9 and Supplementary Figs. 12 and 13).

To unify and simplify large-scale workflows, we developed AFU, a module that integrates the functionalities of AFLP and AFVS into a cohesive interface. AFU enables users to perform ligand preparation and docking within a single automated pipeline, dramatically reducing setup complexity and increasing accessibility for nonspecialists. This streamlined design supports rapid iteration and facilitates integration with other tools in computational drug discovery and artificial

intelligence (AI)-driven design workflows (Supplementary Figs. 3–7 and 14 and Methods).

Enamine REAL Space

Virtual screening performance scales notably with library size, because larger screening sizes tend to yield both higher potencies and improved hit rates (that is, the number of confirmed hits divided by the number of tested compounds)^{4,6–8}. Until recently, the largest publicly available ready-to-dock libraries, including the 2018 version of the REAL Database⁸ and the ZINC20 library¹⁴, each contained approximately 1.5 billion molecules. To enable virtual screens of an order-of-magnitude larger scale, we prepared the 2022 version of the Enamine REAL Space into a ready-to-dock format, comprising 68.7 billion commercially available on-demand small molecules.

The current version of the Enamine REAL Space (2022q1-2, accessed 15 November 2022) contains 31,507,987,117 enumerated molecules (becoming 69 billion after ligand preparation, for example, because of stereoisomer and tautomer enumeration). These compounds are synthetically accessible derivatives of 137,000 established building blocks, assembled using 167 internally validated one-pot synthetic protocols developed at Enamine^{15–22}. This combinatorial framework enables high-throughput compound generation with a reported synthetic success rate of 82%, on the basis of 386,000 experimental reactions conducted in 2021. In total, 99% of the compounds in the library conform to Lipinski's Rule of Five²³. Structurally, it offers exceptional chemical diversity, comprising 1,098,811,629 unique Murcko scaffolds.

Following ligand preparation with AFLP, the number of ready-to-dock molecules expanded to 68.7 billion in the REAL Space. This transformation involved a comprehensive set of preprocessing steps, including stereoisomer and tautomer enumeration, desalting, neutralization, protonation state prediction, 3D conformer generation and conversion into formats compatible with a broad range of docking programs. In addition, 28 physicochemical properties and cheminformatic descriptors were calculated for each compound (Fig. 2 and Supplementary Table 2). To organize this vast chemical space and enable efficient navigation, a subset of 18 key molecular properties were selected to construct an 18-dimensional matrix, in which each axis corresponds to a descriptor such as molecular weight, octanol–water partition coefficient (logP), hydrogen bond donor and acceptor counts, rotatable bond count, topological polar surface area (TPSA), aqueous

solubility (logS), aromatic ring count, molecular refractivity, formal charge, positive and negative charge counts, fraction of sp^3 -hybridized carbons, number of chiral centers, halogen and sulfur atom counts and stereoisomer count. Each dimension was discretized into multiple property intervals, resulting in a combinatorial matrix of 24.4 billion potential tranches (that is, matrix elements). Molecules from the REAL Space occupy approximately 12 million of these tranches, allowing for highly granular, property-based selection of compound subsets. The 18 properties were chosen to balance chemical relevance for drug discovery with computational efficiency, as excessive dimensionality would lead to sparsity and reduce practical tractability (Fig. 2). The binning schemes and occupancy statistics for each dimension are provided in Extended Data Fig. 1.

The REAL Space is not only vast but also chemically rich and structurally diverse, particularly in its representation of 3D and stereochemically complex molecules^{14,24}. Notably, over 62 billion molecules (74%) contain at least one chiral center, contributing to what is often referred to as ‘nonflat’ chemical space, a feature increasingly associated with favorable pharmacological properties^{25,26}. Ligand–receptor binding is inherently 3D and the presence of multiple stereoisomers allows subtle shape complementarity to be exploited. In this context, differences in docking scores across stereoisomers can serve as a proxy for specific target engagement, aiding in the prioritization of high-affinity binders. The REAL Space also includes a substantial number of halogenated compounds, 23 billion in total, including 17.4 billion containing fluorine atoms, which are well suited for ¹⁹F nuclear magnetic resonance (NMR)-based binding assays that do not require isotopically labeled proteins. Despite its breadth, the 18-dimensional matrix remains sparsely populated, revealing large underexplored regions of chemical space. These unoccupied tranches present compelling opportunities for prospective synthesis campaigns to populate additional chemical spaces, enabling the development of molecules to selectively target previously undruggable regions of the proteome.

To facilitate exploration of this high-dimensional chemical library, we developed two interactive web interfaces (Extended Data Fig. 4 and Supplementary Fig. 11), each built with different rendering frameworks to optimize performance across devices and browsers. These interfaces are a projection of the 18-dimensional matrix onto a two-dimensional (2D) matrix, allowing the users to dynamically visualize the properties. Each of these two axes of the 2D projection is based on one of the 18 properties that define the full matrix and the user can select the two properties. Although the visualization is limited to 2D projections, the range of the 18 properties (of the corresponding 18 dimensions) can be selected interactively independently of the visualization. Selected subsets can then be used directly for virtual screens through AFVS. The REAL Space is available in PDB, PDBQT, MOL2 and SDF formats, ensuring compatibility with most structure-based docking programs. Additionally, the library is distributed in SMILES, SELFIES and Parquet formats to support cheminformatics, generative modeling and large-scale analytics. Both the full library and the user-defined subsets can be accessed through the AdaptiveFlow project homepage or through the AWS Open Data repository, enabling direct use in the cloud without the need for local storage.

Adaptive ULVSs

Screening ultralarge ligand libraries containing billions of molecules is computationally intensive, especially when using classical physics-based docking programs. As the number of compounds in the library increases into the tens of billions, the cost of screening the library becomes prohibitively expensive.

Virtual screening is a process of searching a vast chemical space for the best candidate that matches the structural and biophysical properties of the target site in a given receptor. However, the fraction of compounds in an ultralarge screen that will optimally engage

a specific site in the receptor is often very small and share specific molecular characteristics. To be cost and time effective, it is desirable to focus the screen on the optimal chemical space that is likely to contain hits for a given target. To rapidly narrow the search space for a given target, we introduce the ATG-VS technique. This approach consists of the following steps:

1. **Library decomposition.** On average, a tranche in the current library has approximately 5,600 molecules and, for each tranche, up to ten representative ligands are selected for a pilot screen. The number of representatives per tranche (1–10) can be selected by the user. As there are approximately 12 million tranches (that contain molecules from the REAL Space), if one representative per tranche is selected, the total number of molecules screened in the prescreen would be approximately 12 million for the entire library. If more than one representative molecule per tranche is chosen by the user, then these additional molecules are automatically selected to be diverse from each other within the given tranche.
2. **Screening representative molecules (ATG prescreen).** Instead of screening the entire library, a prescreen of only the selected representative ligands is performed. In addition, instead of screening representative ligands from all the 12 million tranches, users can apply an optional dynamic filtering procedure that corresponds to the 18-dimensional properties of the ligands. Here, users can manually apply criteria that the molecules or tranches need to satisfy, for example, on the basis of the properties of the targeted site, which can greatly reduce the number of dockings required in the prescreen. After the prescreen, the subspace (tranches) of the ligand library with the best docking scores among representative ligands is selected for the next screening stage.
3. **Selection of active tranches.** AdaptiveFlow automatically selects the most promising tranches on the basis of the docking scores of the ATG prescreen. The total number of ligands to be screened in the next stage can be specified by the user.
4. **Active cell screening (ATG primary screen).** In the next stage, the tranches that were selected after the prescreen are fully screened (for example, in Fig. 3, the three tranches that contain the best ligands are marked purple). To further refine this selection, an ML classification model (Methods) is used to identify the most promising molecules within each tranche. The model is trained on Morgan fingerprints (size: 1,024) of the prescreened compounds, using the docking scores obtained during the prescreen as labels. It applies a threshold-based classification approach, where the 75th percentile of docking scores serves as the decision boundary. Only molecules classified as high-confidence binders are prioritized for full docking, ensuring that computational resources are allocated efficiently to the most promising candidates.

ATG-VSs greatly reduce the dimension of the chemical space to be screened, thereby reducing the required computation time. AFVS can also prepare new and/or proprietary ligand libraries with ATG-VS support, which requires that the prepared ligand libraries are structured into a multidimensional tranche matrix with preassigned representative molecules for each tranche. An optional third-stage screen can be performed to rescore the best hits with higher accuracy, such as by including receptor flexibility. AdaptiveFlow can carry out all three stages in a fully automated workflow.

AdaptiveFlow is a highly flexible platform that supports funnel-type screening strategies, either as a standalone method or in combination with ATG-VS. Users can perform iterative filtering by screening a subset of the library (or the full library, if computationally feasible) to generate an initial ranking, before rescreening the top candidates with increased exhaustiveness or different scoring functions. This allows for an arbitrary number of screening stages to be deployed.

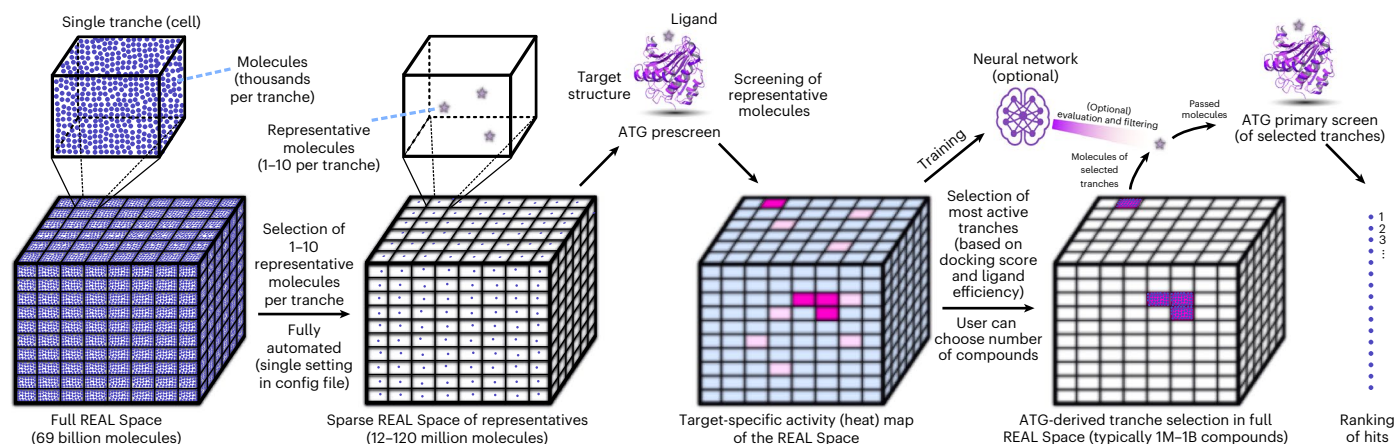


Fig. 3 | Conceptual workflow of ATG-VS. The REAL Space, divided into an 18-dimensional grid, populates over 12 million individual tranches (cells) represented in the figure schematically as a 3D matrix; each of these tranches has on average 5,600 molecules. During the first step of ATG-VS, the AFVS module of AdaptiveFlow selects a few (1–10) representative molecules in each tranche and docks these molecules to the target receptor. This screening mode (and stage)

is referred to as a prescreen and allows AdaptiveFlow to detect the tranches of the library that contain molecules of high potential binding affinity to the given target receptor or binding site. AdaptiveFlow then selects these tranches and screens them fully against the target receptor. In addition, an optional ML model can be trained on the ATG Prescreen results, which can evaluate the molecules of the selected tranches during the ATG Primary Screen.

Additionally, users can manually select library subsets on the basis of molecular properties or generate smaller custom libraries through the AFLP module to facilitate multistage workflows.

To investigate the effectiveness of the various parameters (including the number of representative molecules from each tranche) and the docking exhaustiveness (which is a measure of how extensively the ligand conformational space is sampled), we carried out multiple ATG prescreens on an array of targets. We examined the following scenarios: (1) one representative molecule per tranche with an exhaustiveness of 1; (2) ten representative molecules with an exhaustiveness of 1; and (3) one representative molecule with an exhaustiveness of 10. A higher exhaustiveness parameter results in a more thorough search of the conformational space but more computation time is needed. The results of these trials can be seen in Extended Data Fig. 5. As shown in the figure, the results (as a function of the distribution of the top docking scores and the ranking of the tranches within each property category) are relatively similar for all three scenarios. Therefore, ATG prescreens with an exhaustiveness 1 and one representative molecule per tranche are expected to be sufficient for many applications. In this case, an ATG prescreen involves docking 12 million molecules to explore the library of 69 billion compounds, which is >5,000-fold less computational effort than screening all molecules of the library. If a second-stage screen with higher accuracy is added in the context of multistaging, it may be beneficial to increase the exhaustiveness⁸.

To systematically evaluate the performance of ATG-VS relative to standard ULVSs, we conducted benchmarking experiments across ten diverse protein targets, including kinases, phosphatases, G-protein-coupled receptors and protein–protein interaction interfaces (Supplementary Table 7). These evaluations were performed in two settings: (1) test-based benchmarks using a 5-million-compound subset of chemical space (Fig. 4) and (2) large-scale production runs using the full 69-billion-compound Enamine REAL Space (Extended Data Figs. 6 and 7). The production-scale benchmarks compared standard ULVSs to ATG-VS without the optional active learning component, while the subset-based studies also assessed ATG-VS with active learning enabled. Across these targets, we evaluated performance in the following scenarios:

- (i) Comparison to standard ULVSs (production-scale benchmarks). We evaluated performance at full production scale (Extended Data Figs. 6 and 7) by comparing the top 50 docking scores

obtained from two approaches: (1) a standard ULVS involving 100 million randomly selected compounds from the full Enamine REAL Space and (2) the ATG-VS method without active learning. For ATG-VS, we conducted a prescreen of 12 million compounds, comprising ten representative ligands per tranche across the entire library, followed by a primary screen of up to 1 million compounds. To assess the impact of screening depth, we systematically varied the number of compounds in the primary screen (10,000, 100,000 and 1 million) across ten targets. Active learning was not applied in this setting to isolate and evaluate the impact of tranche prioritization alone; however, on the basis of the test-based benchmarks, we expect that incorporating active learning could yield further improvements in efficiency and enrichment. We found that ATG-VS achieves comparable or superior top docking scores to standard ULVSs while requiring dramatically fewer docking evaluations, highlighting its effectiveness and scalability in production-scale virtual screening campaigns.

- (ii) Impact of ML (test-based benchmarks). To further evaluate the benefit of incorporating an ML classifier into the ATG workflow, we conducted experiments using a subsampled version of the REAL Space comprising 5 million compounds. We compared three approaches: (1) a random screen of 1 million compounds ('1M random' in Fig. 4); (2) the ATG-VS method, which involved decomposing the library based on 18 molecular properties, performing a prescreen and selecting compounds for a primary screen, with a total of 100,000 docking evaluations across both stages ('100K ATG'); and (3) an enhanced ATG-VS approach in which an ML model, trained on docking scores and Morgan fingerprints from the prescreen, was used to prioritize compounds for the primary screen ('100K ATG + AL'). It should be noted that, in the case of the ATG and ATG + AL scenarios, the total number of docking evaluations including the prescreen was 100,000, compared to 1 million in the random screen. As shown in Fig. 4, the ML-augmented ATG-VS consistently improved the enrichment of high-affinity virtual hits relative to the standard ATG-VS. Although both versions of the ATG screen sampled only 100,000 molecules, they outperformed the random screen sampling of 1 million molecules.

While ATG-based screens consistently outperformed random screening, the extent of improvement varied depending on the

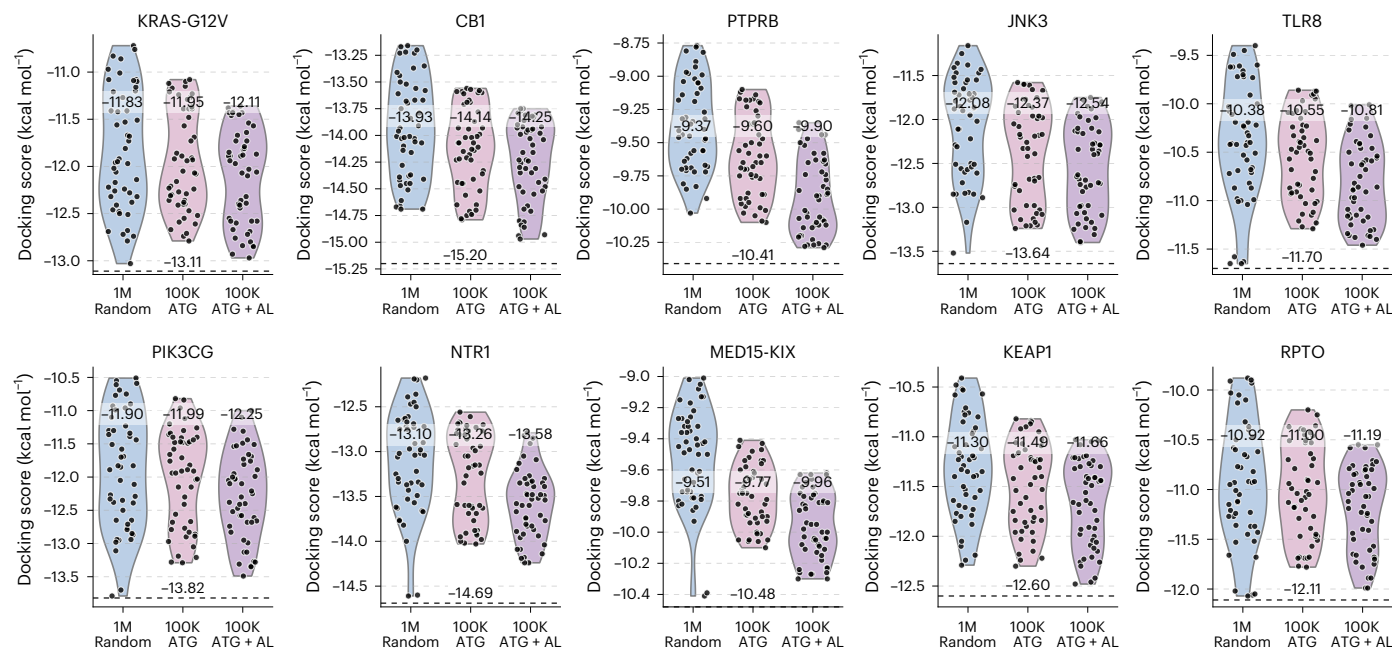


Fig. 4 | Benchmark comparison of ATG-VS and standard ULVS across ten protein targets. Violin plots of the docking scores (the more negative, the better) of the top 50 virtual screening hits obtained using three approaches: (1) standard ULVS of 1 million random compounds (blue); (2) ATG screening with 100,000 molecules (rose); and (3) ATG screening augmented with an additional active learning step (purple). Black dots represent individual data points and the numbers above the violin plots indicate the mean docking scores of the top

50 hits. The horizontal dashed lines indicate the best docking score obtained from a brute-force screen of the entire library, serving as a reference for the absolute best possible hit (5-million-compound subset). For most target proteins, ATG primary screens (100,000 molecules) with or without active learning performed comparably to standard ULVS (1 million random compounds), achieving a tenfold reduction in the number of compounds screened while maintaining similar docking performance.

biophysical properties of the target. Proteins with complex, feature-rich binding pockets, particularly those combining extensive hydrophobic surfaces with well-defined polar subpockets, derived the greatest benefit from ATG-based prioritization and ML enhancement. Taken together, these findings underscore the scalability, efficiency and adaptability of ATG-VS. By intelligently narrowing the search space and optionally integrating ML for compound prioritization, AdaptiveFlow enables high-performance virtual screening at the scale of tens of billions of molecules.

To evaluate the sensitivity of our binning strategy to different sampling algorithms and scoring functions, we repeated the benchmarks using Smina (implementing the Vinardo scoring function) and GWOVina (incorporating flexible-side-chain docking). The physicochemical property-based binning remains effective independent of the docking protocol used. As shown in Supplementary Fig. 13, the ATG-VS approach consistently identified high-scoring compounds with speedups and enrichment factors comparable to those observed with QuickVina 2 (ref. 27). Specifically, the average predicted binding affinity of the top 50 molecules improved relative to standard random screens, while the computational costs remained dramatically lower. This demonstrates that the descriptor-based binning correlates well with binding potential across docking protocols of varying sophistication.

ML and cloud support

AdaptiveFlow introduces expanded capabilities designed to support AI and ML workflows across molecular biology and drug discovery. In addition to integrating state-of-the-art ML-based docking programs, including EquiBind²⁸, DiffDock¹³ and TANKBind¹², the platform offers several features aimed at facilitating the development and application of ML models. The AFLP module computes a wide range of physicochemical properties commonly used as features or optimization targets in ML pipelines, including the quantitative estimate of drug likeness²⁹ and logP³⁰. AFLP also supports the generation of SELFIES molecular representations^{31,32}, which are widely adopted in generative modeling

because of their syntactic robustness and ability to encode valid molecules across arbitrary token sequences³³ (Supplementary Tables 1–4, 8, 9 and 11). Additionally, AFVS introduces numerous deep-learning-based docking approaches by modularly assembling scoring functions and pose prediction tools into comprehensive docking workflows. To further streamline ML-driven workflows, the AFU module provides a fully integrated pipeline for ligand preparation and docking. This unified interface simplifies data generation for model training and enables seamless incorporation of docking scores as objective functions in generative or optimization-based ML models³⁴. With support for approximately 1,500 docking protocols accessible through a single command-line interface, AFU eliminates the need for users to individually configure and learn multiple docking programs, thus greatly lowering the barrier to entry for ML and drug discovery communities alike. Together, these features position AdaptiveFlow as a versatile platform not only for applying ML methods to virtual screening but also for developing ML architectures and pipelines for structure-based drug discovery (Extended Data Figs. 8 and 9).

In parallel with its support for ML, AdaptiveFlow is optimized for efficient deployment in cloud computing environments. The platform runs natively on the popular workload manager Slurm (<https://slurm.schedmd.com>), widely supported across providers such as Google Cloud, Microsoft Azure and AWS Batch (<https://aws.amazon.com/batch/>). Users provide a SLURM submission template (for example, `template1.slurm.sh`) in the configuration. AdaptiveFlow populates this template with the necessary variables (job arrays, partitions and memory) at runtime. This decouples the workflow logic from the specific cluster configuration. Users can adapt the submission script to their local requirements (for example, specific headers, modules or partition names) without modifying the core Python codebase. Furthermore, the size of computing jobs (both size in cores and length of time) can be configured according to site-specific cluster configurations.

In this study, we demonstrate near-linear scalability on AWS using over 5.6 million vCPUs, advancing parallelization in cloud-based drug

discovery. This level of scalability enables the completion of ULVSS involving billions of compounds within hours, timelines that would otherwise require weeks or months on traditional on-premise compute clusters. By combining cloud-native performance with ML-native features, AdaptiveFlow serves as a high-performance and accessible foundation for next-generation AI-driven drug discovery.

Experimental validation of AdaptiveFlow

To experimentally validate the capabilities and performance of AdaptiveFlow, we selected two therapeutically important target proteins, FSP1 and PARP1. FSP1 is a recently characterized oxidoreductase that has a central role in protecting cells from ferroptosis, an iron-dependent form of regulated cell death triggered by membrane lipid peroxidation^{35–38}. FSP1 maintains pools of reduced coenzyme Q₁₀ and vitamin K hydroquinone using NAD(P)H, operating independently of GPX4 to suppress oxidative membrane damage. Inhibiting FSP1 disrupts this defense, sensitizing cells, especially redox-imbalanced cancer cells³⁹, to ferroptosis. Despite its importance, only a few FSP1 inhibitors have been described and their binding to FSP1 has been structurally characterized^{10,40–42}, making it an ideal target for exploring novel inhibitor space using AdaptiveFlow. We also targeted PARP1, a well-established target involved in DNA single-strand break repair, for benchmarking AdaptiveFlow on a mechanistically distinct, clinically validated enzyme class.

Structural studies of FSP1

To streamline hit discovery and structural studies, we applied ancestral sequence reconstruction, a method known to yield highly stable proteins while preserving functional and structural fidelity to modern homologs^{43,44}. Using this approach, we reconstructed the tetrapod ancestor of human FSP1, which shares 72.5% sequence identity and retains all active-site residues (Supplementary Figs. 15 and 16). The ancestral FSP1 variant expressed efficiently in *Escherichia coli* (~100 mg l⁻¹ culture), bound FAD, exhibited high thermostability (T_m = 62 °C; Supplementary Fig. 16b) and retained full enzymatic activity (Supplementary Table 10 and Supplementary Figs. 17 and 18). Crucially, it yielded high-quality, well-diffracting (Supplementary Table 11), enabling detailed structural analyses.

We determined high-resolution crystal structures of FSP1 bound to FAD, bound to FAD and NAD⁺ and in complex with FAD, NAD⁺ and coenzyme Q₁ (Extended Data Fig. 10, Supplementary Fig. 19a and Supplementary Table 11). The overall structure of the ancestral protein is very similar to that of the human enzyme (Supplementary Fig. 19b). NAD⁺ triggers a 2-Å shift of the flavin ring and its binding mode is fully consistent with a hydride transfer mechanism for flavin reduction as the nicotinamide C4 atom is located right in front of the isoalloxazine N5 (Extended Data Fig. 10b). Coenzyme Q₁ binds between two aromatic side chains of Phe354 and Y290, with its quinone moiety at the edge of the pyrimidine ring of the flavin in a canyon leading from the protein surface (Extended Data Fig. 10c,d). Specifically, quinone oxygen engages the N(3)H group of the flavin through a hydrogen bond. This geometry and the proximity between the two molecules can enable the facile electron transfer from the flavin to coenzyme Q. Comparison between the FSP1–FAD–NAD⁺ and FSP1–FAD–NAD⁺–coenzyme Q₁ complexes shows that NAD⁺ binding is unaffected by the presence of coenzyme Q. Thus, it is demonstrated that coenzyme Q and NAD(P)⁺ can simultaneously bind to the protein and this structural attribute perfectly aligns with the ternary complex mechanism evidenced by the kinetics data (Supplementary Fig. 18). The structure also suggests that there is ample space for the oxygen to access the flavin N5 locus where its reaction with the flavin is known to take place. This feature explains why FSP1 can also operate as a NAD(P)H oxidase⁴⁵. This detailed experimental visualization of the NAD(P)H and quinone-binding modes and their interactions provided the foundation for inhibitor discovery through in silico screenings targeting the substrate-binding sites.

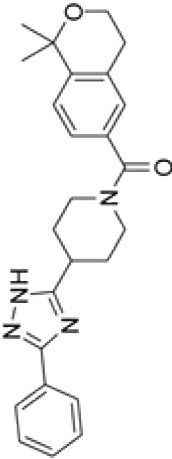
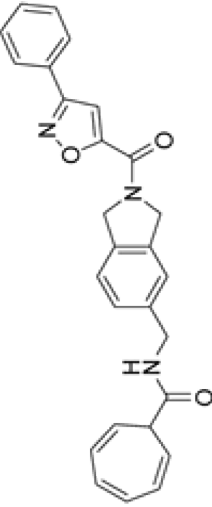
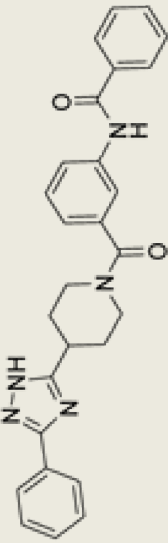
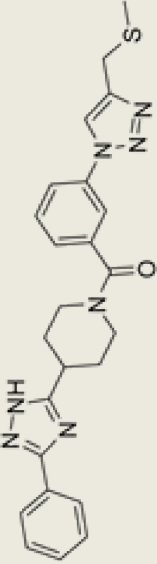
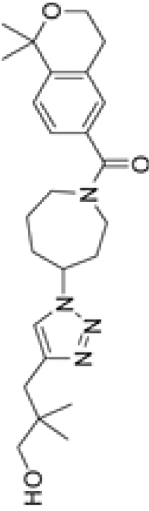
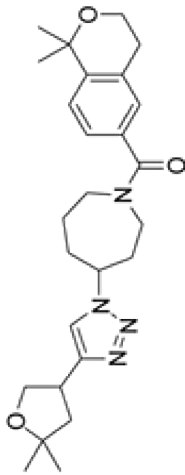
Targeting FSP1 with AdaptiveFlow

To assess the effectiveness of AdaptiveFlow, we screened FSP1 inhibitors targeting its coenzyme Q site using AdaptiveFlow. Our high-resolution crystal structure of FSP1 in complex with FAD and NAD⁺ (PDB 9IFT) was used as a receptor. Maestro (Schrödinger) was applied to prepare the receptor's structure for screening, by removing solvent molecules, assigning missing hydrogen bonds, correcting protonation states at pH 7.4 and assigning correct bond orders. AutoDockTools (version 1.5.7)⁴⁶ was used to convert the prepared structure from the PDB format to the PDBQT format. The coenzyme Q site of FSP1 (with FAD and NAD⁺ bound to the corresponding sites) was targeted. In the ATG Prescreen, one representative compound was used per tranche (resulting in 12 million dockings) and, on the basis of these results, 10 million compounds were chosen for the primary screen using the ATG-VS method (without ML feature). Basic filters were applied to the top hit candidates, such as logP < 5, logS > -5, TPSA < 140 Å², number of stereo centers ≤ 3 and excluding compounds with problematic functional groups and predicted toxicity (mutagenic compounds, compounds with reproductive effects or tumorigenic compounds). In total, 42 compounds from 37 different clusters with zero alerts suggested by SwissADME⁴⁷ were ordered from Enamine and 33 compounds were synthesized and experimentally tested.

Experimental validation: inhibitors extend from the coenzyme Q site to the NAD site

We assessed the potential binding and inhibitory activities of the synthesized 33 virtual hits from AdaptiveFlow, using thermal shift and NADH depletion assays. Two of the hits elicited stabilizing shifts of at least 1.5 °C in the protein melting temperature and inhibited FSP1 oxidoreductase activity by at least 50% at a 10 μM concentration. These two compounds were found to potentially inhibit coenzyme Q-reduction activity (Table 1 and Supplementary Fig. 20). Their inhibition constant (K_i) values are comparable to or better than those reported for previously described FSP1 inhibitors (0.283 μM for afi-FSP1-1 and 0.777 μM for afi-FSP1-2; Supplementary Table 12 and Supplementary Figs. 21 and 22)^{10,40–42}. Critically, a cellular thermal shift assay⁴⁸ performed on FSP1-expressing HEK293 cells showed that afi-FSP1-1 and afi-FSP1-2 effectively engage human FSP1 in cells, as evidenced by thermal shifts of 2.8 °C and 3.1 °C, respectively (Supplementary Fig. 23a–j). The activity of afi-FSP1-1 was tested against a T cell acute lymphoblastic leukemia cell line (Jurkat) that expresses FSP1 (Supplementary Fig. 23m) and is known to undergo ferroptosis⁴⁹. Cells proved to be sensitive to the combination of afi-FSP1-1 with 0.5 μM erastin, a ferroptosis inducer. The activity of afi-FSP1-1 matched that shown by the widely used iFSP1 (ref. 10). When used as a single agent, afi-FSP1-1 showed no cell toxicity at concentrations up to 10 μM. The proferroptotic effect induced by 25 μM inhibitor was completely rescued by liproxstatin-1, a ferroptosis inhibitor (Supplementary Fig. 23n–p). We determined the crystal structure of the most potent compound (afi-FSP1-1) bound to FSP1–FAD–NAD⁺, which reveals that the inhibitor occupies the coenzyme-Q-binding site in a mode consistent with the predicted docking pose (Fig. 5a,b and Supplementary Fig. 19a). With its elongated conformation, afi-FSP1-1 extensively interacts with the protein by contacting several aromatic side chains, including Y290 and F354, residues that sandwich the quinone ring of coenzyme Q. Y290 is involved in a π -stacking interaction with the benzene ring of the isochroman moiety, whereas F354 and F15 contribute to the binding through hydrophobic interactions with the piperidine and a π -stacking interaction with the triazole group. The central carbonyl oxygen of afi-FSP1-1 forms an H-bond interaction with the N(3)H group atom of the flavin akin to the interaction established by the carboxyl oxygen of the coenzyme Q's ubiquinone head (Fig. 5b and Extended Data Fig. 10c). Furthermore, the inhibitor's central carbonyl oxygen consistently interacts with the backbone amine group (NH) of L323 through a water bridge (Supplementary Fig. 24). Despite its proximity with the inhibitor, the binding conformation of NAD⁺ is

Table 1 | Competitive K_i and binding K_d values of the most potent FSP1 inhibitors

	K_i	K_d		K_i	K_d
afi-FSP1-1	$0.283^a \pm 0.05$	0.098	afi-FSP1-2	0.777 ± 0.22	0.266
					
afi-FSP1-3	0.520 ± 0.11	0.184	afi-FSP1-4	0.284 ± 0.07	0.198
					
afi-FSP1-5	1.000 ± 0.32	n.d.	afi-FSP1-6	3.940 ± 0.71	n.d.
					

K_i values were determined by monitoring NADH oxidation at 340nm (pH 7.2) using variable NADH concentrations (5–200 μ M), 100 μ M coenzyme Q_1 and 0.1–0.2 μ M FSP1 (Supplementary Fig. 20). Dissociation constant (K_d) values were derived from independent binding experiments based on surface plasmon resonance (Supplementary Fig. 22). The experiments were performed using the racemic mixture of compound afi-FSP1-5 and the diastereomeric mixture of compound afi-FSP1-6. Inhibitor concentrations ranged from micromolar to low nanomolar. Supplementary Table 12 compares the K_i values of the afi-FSP1 compounds with those of known FSP1 inhibitors, whereas the IC_{50} values measured using the indirect resazurin assay, commonly used in previous studies^{10,42}, are listed in Supplementary Table 16. n.d., not determined. ^a The K_i measured by varying coenzyme Q_1 (5–100 μ M) at fixed 100 μ M NADH is $0.243 \pm 0.06 \mu$ M.

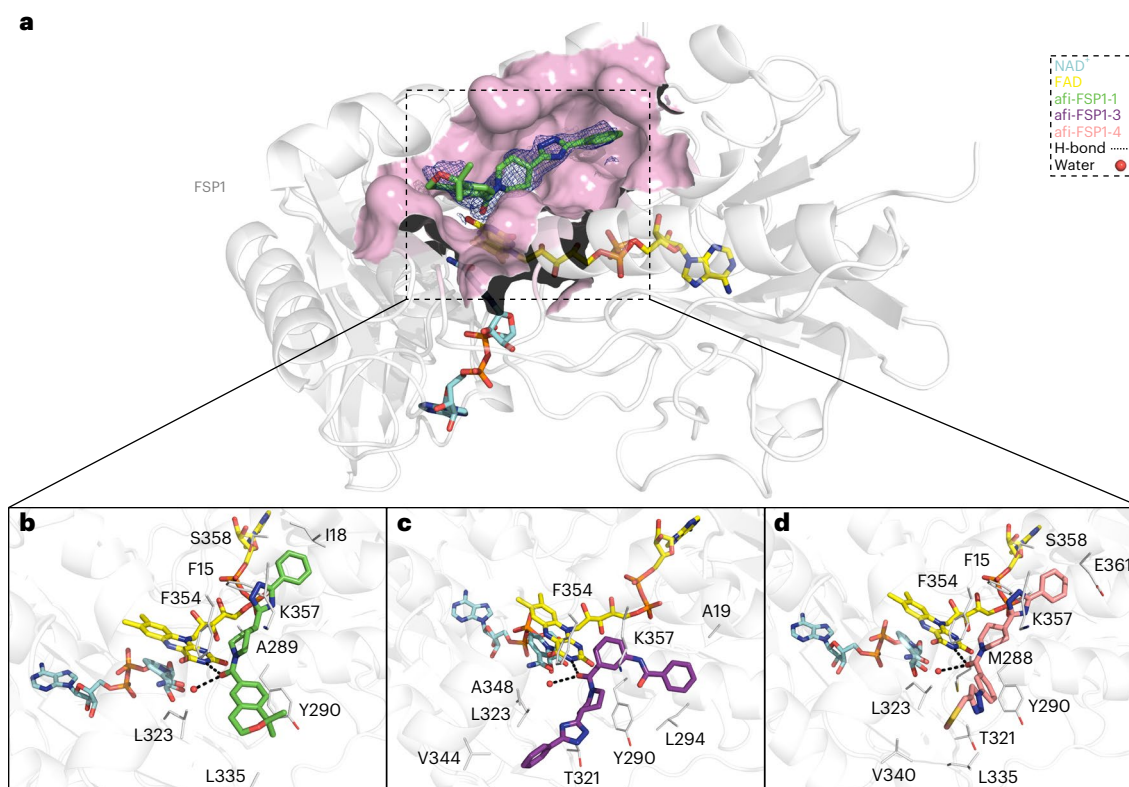


Fig. 5 | Co-crystal structures of inhibitors binding to FSP1. a, Overall structure of FSP1 in complex with inhibitor ari-FSP1-1. The polder omit map for the inhibitor is contoured at the 3.0 σ level. **b–d**, Close-up views of the binding site with

ari-FSP1-1 (**b**; green), ari-FSP1-3 (**c**; violet) and ari-FSP1-4 (**d**; pink) bound in the region normally occupied by coenzyme Q₁ adjacent to FAD (yellow) and NAD⁺ (cyan). All inhibitors displace coenzyme Q₁ and interact directly with the catalytic cavity.

identical to that observed both in the absence of a ligand and in the complex with coenzyme Q₁. Similar features characterize the binding compound ari-FSP1-2 whose bulkier structure enables even more extended interaction with the protein (Supplementary Figs. 24 and 25). Expanding on these results, we conducted a structure–activity relationship study. The 3D structure of the FSP1–ari-FSP1-1 complex was used as input for an additional round of *in silico* screening. This procedure returned 20 analogs of compound ari-FSP1-1 that were experimentally analyzed (Supplementary Fig. 26). Rewardingly, four additional molecules (ari-FSP1-3–ari-FSP1-6) displayed inhibitory activities with K_i values in the three-digit nanomolar range for achiral molecules to low micromolar range for molecules tested as either a racemic mixture or a diastereomeric mixture (Supplementary Figs. 20c–f and 23). Notably, these compounds share an aromatic-amide-saturated heterocycle core as displayed also by ari-FSP1-1 and ari-FSP1-2 (Table 1). Consistent with this common feature, the crystal structures of ari-FSP1-3 and ari-FSP1-4 bound to FSP1–FAD–NAD⁺ show that the ligands occupy the coenzyme Q site with their amide carbonyl oxygen engaged by the flavin N(3)H group through a hydrogen bond and further interacting through a water molecule with the backbone amine group of L323 in FSP1 as observed for ari-FSP1-1 (Fig. 5b–d and Supplementary Figs. 24 and 25). Although compound ari-FSP1-4 replaces the dimethyltetrahydropyran portion of the isochroman group of ari-FSP1-1 with a methylthiomethyltriazole, the two inhibitors bind in a very similar way with their common phenyltriazole–piperidine moiety interacting with F15 and F354. Conversely, the binding of ari-FSP1-3 is more distinct as it adopts a less extended, U-shaped conformation laying on Y290. These binding modes were validated by mutagenesis experiments (Supplementary Tables 13 and 14 and Supplementary Figs. 24 and 27) and correlate with the measured K_i values. Compounds ari-FSP1-1 and ari-FSP1-4 exhibit practically identical K_i values, whereas compound ari-FSP1-3 displays a twofold decrease in potency. Nonetheless,

absorption, distribution, metabolism and excretion (ADME) assays indicated that ari-FSP-3 is the most promising candidate for further medicinal chemistry optimization (Supplementary Information). In addition to being the best inhibitors, compounds ari-FSP1-1, ari-FSP1-2 and ari-FSP1-3 exhibited 10–12-fold reduced efficacy against the NADH oxidase activity of FSP1, while compound ari-FSP1-4 exhibited no efficacy at all (Supplementary Figs. 28 and 29 and Supplementary Table 15). The weaker inhibition reflects the retention of NAD(P)H binding in the inhibited protein as shown by crystal structures and indicates that reactivity with oxygen is only mildly or not affected by inhibitor binding in the coenzyme Q site. We posit that preferentially inhibiting the quinone-reducing activity while sparing the reactivity with oxygen might be pharmacologically beneficial. Inhibitors of this type may convert the enzyme from being antiapoptotic (coenzyme-Q-reducing activity) to being proapoptotic (uninhibited lipid-damaging ROS production). Future work is needed to explore and validate this concept in FSP1 pharmacology.

Applying AdaptiveFlow to PARP1

To further validate the effectiveness of the AdaptiveFlow platform, we selected human PARP1 as a benchmark target because of its well-characterized structure, critical role in DNA repair and relevance in cancer therapy. Following the ATG prescreen and a primary screen of 100 million molecules, 160 promising candidates were synthesized for further experimental verification (Methods). A two-concentration enzymatic HTS identified seven PARP1 inhibitors, four of which (iParp1, iParp2, iParp3 and iParp4) showed sub-250 nM half-maximal inhibitory concentration (IC₅₀) values and were selected for further characterization. Structural similarity analyses confirmed that several hits shared key scaffolds with known PARP1 inhibitors, while others showed different scaffold variations. Remarkably, iParp1 exhibited an IC₅₀ of 8.8 nM, comparable to the US Food and Drug Administration

(FDA)-approved drug olaparib. This result is particularly notable given that our hits were derived from a purely virtual screen without any prior knowledge of known inhibitors or medicinal chemistry optimization of the hit. Protein NMR confirmed direct binding to the active site of PARP1, with chemical shift perturbations supporting the predicted docking poses. X-ray crystallography revealed the 2.05-Å structure of the PARP1 catalytic domain bound to a hydrolyzed form of iParp1, validating its binding mode and key interactions. Clonogenic assays in BRCA1-deficient triple-negative breast cancer cells demonstrated selective cytotoxicity for iParp1 and iParp2, although reduced cellular activity, especially for iParp1, was attributed to poor membrane permeability because of hydrolysis. Details can be found in the Supplementary Information. Overall, this validation exemplifies AdaptiveFlow's ability to efficiently and accurately discover drug-like molecules with clinically relevant potency, streamlining early-stage drug discovery (Supplementary Figs. 30–36 and Supplementary Table 20).

Discussion

The accessible chemical space for drug discovery is expanding at an unprecedented pace, with the number of purchasable, on-demand molecules growing from approximately 2 billion in 2020 to trillions of molecules. Yet, even this remarkable growth represents only a minuscule fraction of the estimated 10^{60} drug-like molecules that populate theoretical chemical space². Identifying high-quality hits within this vast and ever-expanding landscape remains a central challenge in early-stage drug discovery. To address this, we curated the largest ready-to-dock library to date, the Enamine REAL Space, comprising 69 billion commercially accessible small molecules, fully preprocessed and made available in multiple formats for structure-based virtual screening.

While ultralarge libraries increase the likelihood of identifying potent binders and enhance true-hit rates^{4,6–8,50,51}, exhaustively screening billions of compounds impose substantial computational costs, even with today's cloud-based computing capabilities. AdaptiveFlow addresses this challenge through ATG-VS, a highly efficient prescreening strategy that narrows docking efforts to target-specific, chemically diverse subspaces enriched for high-affinity candidates. When applied to the 69-billion-compound REAL Space, ATG-VS reduces screening costs up to 1,000-fold compared to exhaustive searches while preserving strong enrichment of top hits. Additionally, AdaptiveFlow incorporates deep-learning-based docking and GPU acceleration, achieving a further 10–100-fold increase in throughput^{52,53}. To experimentally validate the practical utility of AdaptiveFlow, we focused on two biologically and clinically relevant targets. First, we targeted FSP1, a redox enzyme involved in lipid peroxide detoxification and resistance to ferroptosis and identified a series of nanomolar inhibitors. Structural characterization of FSP1 in complex with cofactors and ligands revealed that both NAD(P)H and coenzyme Q can simultaneously bind in catalytically competent orientations, providing mechanistic insights into its function. The identified hits were validated through crystallography and biochemical assays. We also applied AdaptiveFlow to PARP1, a clinically validated oncology target central to DNA damage repair. Our ultralarge screen yielded seven inhibitors, including iParp1, which displayed an IC_{50} comparable to that of the FDA-approved drug olaparib⁵⁴. These results underscore the power of AdaptiveFlow in rapidly identifying potent hits across diverse protein classes and demonstrate the advantage of adaptive screening over brute-force enumeration. Notably, while our previous identification of nanomolar KEAP1–NRF2 inhibitors⁸ required an unbiased screen of 1.3 billion molecules, the ATG-guided screen for FSP1 achieved comparable success with just 22 million docking instances. Although the effectiveness of chemical space focusing is target dependent, our benchmarking and experimental studies highlight the potential of ATG-VS to tailor the search space to specific binding pockets, greatly enhancing screening efficiency. As the chemical space continues to expand, focusing on target-specific subspaces will no longer be optional, but essential.

With the infrastructure now in place, adding molecules to the library and directing them to their respective tranches is a computational process that scales linearly. As new tranches beyond the 12 million currently populated are filled, only representatives from these newly populated tranches need to be incorporated into the prescreen set. The open-source nature of AdaptiveFlow enables seamless adaptation, extension and development of the platform by both academic and industry users. With the increasing availability of high-resolution structural data, enabled by techniques such as cryo-electron microscopy and de novo structure prediction, combined with site-specific functional insights from gene editing technologies, in silico screening is poised to deliver chemical tools capable of modulating precise biomolecular functions. As modalities like molecular glues and targeted protein degraders mature, the ability to identify small molecules that bind specific protein surfaces with functional consequences is becoming increasingly valuable. AdaptiveFlow's modular architecture supports over 1,500 docking protocols, encompassing combinations of scoring functions, sampling algorithms and ML-based pose predictors such as DiffDock and TANKBind. The platform also natively supports SELFIES representations, physicochemical property calculations and active learning-driven workflows, making it fully compatible with generative models and reinforcement learning pipelines^{55–57}. As generative models continue to advance, platforms such as AdaptiveFlow will be essential for rapidly evaluating designed molecules against biologically relevant structural targets. In conclusion, AdaptiveFlow offers a unified, open-source and massively scalable solution for ULVSS. It provides seamless integration of conventional and AI-driven docking protocols, access to the largest ready-to-dock chemical library and adaptive strategies for efficient navigation of the trillion-scale chemical space.

Online content

Any methods, additional references, Nature Portfolio reporting summaries, source data, extended data, supplementary information, acknowledgements, peer review information; details of author contributions and competing interests; and statements of data and code availability are available at <https://doi.org/10.1038/s41587-026-03217-x>.

References

1. Mennen, S. M. et al. The evolution of high-throughput experimentation in pharmaceutical development and perspectives on the future. *Org. Process Res. Dev.* **23**, 1213–1242 (2019).
2. Bohacek, R. S., McMartin, C. & Guida, W. C. The art and practice of structure-based drug design: a molecular modeling perspective. *Med. Res. Rev.* **16**, 3–50 (1996).
3. Ellingson, S. R., Dakshanamurthy, S., Brown, M., Smith, J. C. & Baudry, J. Accelerating virtual high-throughput ligand docking: current technology and case study on a petascale supercomputer. *Concurr. Comput.* **26**, 1268–1277 (2014).
4. Lyu, J. et al. Ultra-large library docking for discovering new chemotypes. *Nature* **566**, 224–229 (2019).
5. Stein, R. M. et al. Virtual discovery of melatonin receptor ligands to modulate circadian rhythms. *Nature* **579**, 609–614 (2020).
6. Alon, A. et al. Structures of the σ_2 receptor enable docking for bioactive ligand discovery. *Nature* **600**, 759–764 (2021).
7. Gorgulla, C. *Free Energy Methods Involving Quantum Physics, Path Integrals, and Virtual Screenings: Development, Implementation and Application in Drug Discovery*. PhD thesis, Freie Universität Berlin (2018).
8. Gorgulla, C. et al. An open-source drug discovery platform enables ultra-large virtual screens. *Nature* **580**, 663–668 (2020).
9. Grygorenko, O. O. et al. Generating multibillion chemical space of readily accessible screening compounds. *iScience* **23**, 101681 (2020).

10. Doll, S. et al. Fsp1 is a glutathione-independent ferroptosis suppressor. *Nature* **575**, 693–698 (2019).
11. Lord, C. J. & Ashworth, A. PARP inhibitors: synthetic lethality in the clinic. *Science* **355**, 1152–1158 (2017).
12. Lu, W. et al. TANKBind: trigonometry-aware neural networks for drug–protein binding structure prediction. Preprint at *bioRxiv* <https://doi.org/10.1101/2022.06.06.495043> (2022).
13. Corso, G., Stärk, H., Jing, B., Barzilay, R. & Jaakkola, T. DiffDock: Diffusion steps, twists, and turns for molecular docking. Preprint at <https://doi.org/10.48550/arXiv.2210.01776> (2022).
14. Irwin, J. J. et al. ZINC20—a free ultralargescale chemical database for ligand discovery. *J. Chem. Inf. Model.* **60**, 6065–6073 (2020).
15. Bogolubsky, A. V. et al. Bis(2,2,2-trifluoroethyl) carbonate as a condensing agent in one-pot parallel synthesis of unsymmetrical aliphatic ureas. *ACS Comb. Sci.* **16**, 303–308 (2014).
16. Bogolubsky, A. V. et al. Sulfonyl fluorides as alternative to sulfonyl chlorides in parallel synthesis of aliphatic sulfonamides. *ACS Comb. Sci.* **16**, 192–197 (2014).
17. Bogolubsky, A. V. et al. A one-pot parallel reductive amination of aldehydes with heteroaromatic amines. *ACS Comb. Sci.* **16**, 375–380 (2014).
18. Bogolubsky, A. V. et al. One-pot parallel synthesis of alkyl sulfides, sulfoxides, and sulfones. *ACS Comb. Sci.* **17**, 348–354 (2015).
19. Bogolubsky, A. V. et al. 2,2,2-Trifluoroethyl chlorooxoacetate—universal reagent for one-pot parallel synthesis of N^1 -aryl- N^2 -alkyl-substituted oxamides. *ACS Comb. Sci.* **17**, 615–622 (2015).
20. Tolmachev, A. et al. Expanding synthesizable space of disubstituted 1,2,4-oxadiazoles. *ACS Comb. Sci.* **18**, 616–624 (2016).
21. Bogolubsky, A. V. et al. An old story in the parallel synthesis world: an approach to hydantoin libraries. *ACS Comb. Sci.* **20**, 35–43 (2018).
22. Savych, O. et al. One-pot parallel synthesis of 5-(dialkylamino) tetrazoles. *ACS Comb. Sci.* **21**, 635–642 (2019).
23. Lipinski, C. A., Lombardo, F., Dominy, B. W. & Feeney, P. J. Experimental and computational approaches to estimate solubility and permeability in drug discovery and development settings. *Adv. Drug Deliv. Rev.* **23**, 3–25 (1997).
24. Tomberg, A. & Boström, J. Can easy chemistry produce complex, diverse, and novel molecules? *Drug Discov. Today* **25**, 2174–2181 (2020).
25. Lovering, F., Bikker, J. & Humblet, C. Escape from flatland: increasing saturation as an approach to improving clinical success. *J. Med. Chem.* **52**, 6752–6756 (2009).
26. Lovering, F. Escape from flatland 2: complexity and promiscuity. *MedChemComm* **4**, 515–519 (2013).
27. Alhossary, A., Handoko, S. D., Mu, Y. & Kwok, C.-K. Fast, accurate, and reliable molecular docking with QuickVina 2. *Bioinformatics* **31**, 2214–2216 (2015).
28. Stärk, H., Ganea, O., Pattanaik, L., Barzilay, R. & Jaakkola, T. Equibind: geometric deep learning for drug binding structure prediction. In *Proc. of the 39th International Conference on Machine Learning* (eds Chaudhuri, K. et al.) (PMLR, 2022).
29. Bickerton, G. R., Paolini, G. V., Besnard, J., Muresan, S. & Hopkins, A. L. Quantifying the chemical beauty of drugs. *Nat. Chem.* **4**, 90–98 (2012).
30. Wildman, S. A. & Crippen, G. M. Prediction of physicochemical parameters by atomic contributions. *J. Chem. Inf. Comput. Sci.* **39**, 868–873 (1999).
31. Krenn, M., Häse, F., Nigam, A., Friederich, P. & Aspuru-Guzik, A. Self-referencing embedded strings (SELFIES): a 100% robust molecular string representation. *Mach. Learn. Sci. Technol.* **1**, 045024 (2020).
32. Krenn, M. et al. Selfies and the future of molecular string representations. *Patterns* **3**, 100588 (2022).
33. Flam-Shepherd, D., Zhu, K. & Aspuru-Guzik, A. Language models can learn complex molecular distributions. *Nat. Commun.* **13**, 1–10 (2022).
34. Nigam, A. K. et al. Tartarus: a benchmarking platform for realistic and practical inverse molecular design. In *Proc. of the 36th Conference on Advances in Neural Information Processing Systems* (eds Oh, A. et al.) (NeurIPS, 2024).
35. Brown, A. R., Hirschhorn, T. & Stockwell, B. R. Ferroptosis—disease perils and therapeutic promise. *Science* **386**, 848–849 (2024).
36. Lee, J., Yoo, I., Kim, M., Kim, H. & Lee, N. CoQ10 oxidoreductases in ferroptosis and cancer: redox regulation and therapeutic opportunities. *Exp. Mol. Med.* **58**, 1699–1710 (2026).
37. Lange, M. et al. FSP1-mediated lipid droplet quality control prevents neutral lipid peroxidation and ferroptosis. *Nat. Cell Biol.* **27**, 1902–1913 (2025).
38. Mishima, E. et al. A non-canonical vitamin K cycle is a potent ferroptosis suppressor. *Nature* **608**, 778–783 (2022).
39. Sies, H., Mailloux, R. J. & Jakob, U. Fundamentals of redox regulation in biology. *Nat. Rev. Mol. Cell Biol.* **25**, 701–719 (2024).
40. Hendricks, J. M. et al. Identification of structurally diverse FSP1 inhibitors that sensitize cancer cells to ferroptosis. *Cell Chem. Biol.* **30**, 1090–1103 (2023).
41. Nakamura, T. et al. Integrated chemical and genetic screens unveil FSP1 mechanisms of ferroptosis regulation. *Nat. Struct. Mol. Biol.* **30**, 1806–1815 (2023).
42. Zhang, S. et al. Cocrystal structure reveals the mechanism of FSP1 inhibition by FSEN1. *Proc. Natl Acad. Sci. USA* **122**, e2505197122 (2025).
43. Thomson, R. E. S., Carrera-Pacheco, S. E. & Gillam, E. M. J. Engineering functional thermostable proteins using ancestral sequence reconstruction. *J. Biol. Chem.* **298**, 102435 (2022).
44. Nicoll, C. R., Massari, M., Fraaije, M. W., Mascotti, M. L. & Mattevi, A. Impact of ancestral sequence reconstruction on mechanistic and structural enzymology. *Curr. Opin. Struct. Biol.* **82**, 102669 (2023).
45. Guerriere, T. B. et al. Dehydrogenase versus oxidase function: the interplay between substrate binding and flavin microenvironment. *ACS Catal.* **15**, 1046–1060 (2025).
46. Morris, G. M., Huey, R. & Olson, A. J. Using AutoDock for ligand–receptor docking. *Curr. Protoc. Bioinformatics* **Chapter 8**, Unit 8.14 (2008).
47. Daina, A., Michielin, O. & Zoete, V. Swissadme: a free web tool to evaluate pharmacokinetics, drug-likeness and medicinal chemistry friendliness of small molecules. *Sci. Rep.* **7**, 42717 (2017).
48. Molina, D. M. et al. Monitoring drug target engagement in cells and tissues using the cellular thermal shift assay. *Science* **341**, 84–87 (2013).
49. Soula, M. et al. Metabolic determinants of cancer cell sensitivity to canonical ferroptosis inducers. *Nat. Chem. Biol.* **16**, 1351–1360 (2020).
50. Gloriam, D. E. Bigger is better in virtual drug screens. *Nature* **566**, 193–194 (2019).
51. Nigam, A. K. et al. Drug discovery in low data regimes: leveraging a computational pipeline for the discovery of novel SARS-CoV-2 Nsp14-MTase inhibitors. Preprint at *bioRxiv* <https://doi.org/10.1101/2023.10.03.560722> (2023).
52. Santos-Martins, D. et al. Accelerating AutoDock4 with GPUs and gradient-based local search. *J. Chem. Theory Comput.* **17**, 1060–1073 (2021).
53. Tang, S. et al. Accelerating AutoDock Vina with GPUs. *Molecules* **27**, 3041 (2022).

54. Valabrega, G., Scotto, G., Tuninetti, V., Pani, A. & Scaglione, F. Differences in PARP inhibitors for the treatment of ovarian cancer: mechanisms of action, pharmacology, safety, and efficacy. *Int. J. Mol. Sci.* **22**, 4203 (2021).
55. Sanchez-Lengeling, B. & Aspuru-Guzik, A. Inverse molecular design using machine learning: generative models for matter engineering. *Science* **361**, 360–365 (2018).
56. O’Boyle, N. M. et al. Open Babel: an open chemical toolbox. *J. Cheminform.* **3**, 33 (2011).
57. Weininger, D. SMILES, a chemical language and information system. 1. Introduction to methodology and encoding rules. *J. Chem. Inf. Comput. Sci.* **28**, 31–36 (1988).

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

Open Access This article is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License, which permits any non-commercial use, sharing, 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 you modified the licensed material. You do not have permission under this licence to share adapted material derived from this article or parts of it. 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-nc-nd/4.0/>.

© The Author(s) 2026

Domiziana Cecchini ^{1,45}, **AkshatKumar Nigam** ^{2,45}, **Ming Tang** ^{3,4,5,45}, **Joana Reis** ^{6,45}, **Matt Koop**⁷, **Andrea Gottinger** ¹, **Callum Robert Nicoll** ¹, **Yao Wang**⁸, **Abhilash Jayaraj** ^{6,9}, **Süleyman Selim Çınaroglu** ^{10,11}, **Ricarda Törner** ^{6,12}, **Yehor Malets** ¹³, **Minko Gehev**¹⁴, **Krishna M. Padmanabha Das**^{5,9}, **Kelly Churion**⁵, **Jongwan Kim**⁵, **Nidhin Thomas**⁵, **Yong Li**¹⁵, **Hyuk-Soo Seo** ¹⁶, **Sirano Dhe-Paganon** ¹⁶, **Christopher Secker** ^{17,18}, **Mohammad Haddadnia** ^{6,9,19,20,21,22}, **Alexander Hasson** ^{23,24}, **Minkai Li**²⁵, **Abhishek Kumar** ^{5,26}, **Roni Levin-Konigsberg** ^{2,27}, **Eun-Bee Choi**²⁸, **Geoffrey I. Shapiro** ²⁸, **Huel Cox 3rd** ⁶, **Luke Sebastian**⁶, **Chelsea Braithwaite** ⁶, **Puspalata Bashyal**⁶, **Dmytro S. Radchenko** ²⁹, **Aditya Kumar**³⁰, **Lei Yang** ¹⁵, **Pierre-Yves Aquilanti**^{7,31}, **Henry Gabb** ³², **Amr Alhossary** ^{33,34}, **Eric O’Neill**²⁴, **Gerhard Wagner** ⁹, **Alán Aspuru-Guzik**^{20,21,22,35,36,37,38,39,40}, **Yurii S. Moroz**^{29,41,42}, **Charalampos G. Kalodimos**⁵, **Konstantin Fackeldey**^{17,30}, **John D. Schuetz**⁸, **Andrea Mattevi** ¹✉, **Haribabu Arthanari** ^{6,9,43,44}✉ & **Christoph Gorgulla** ^{5,9}✉

¹Department of Biology and Biotechnology, University of Pavia, Pavia, Italy. ²Department of Computer Science, Stanford University, Stanford, CA, USA. ³Frazer Institute, Faculty of Medicine at the Translational Research Institute Australia, The University of Queensland, Brisbane, Queensland, Australia. ⁴School of Biomedical Sciences, Centre for Genomics and Personalised Health at the Translational Research Institute Australia, Queensland University of Technology, Brisbane, Queensland, Australia. ⁵Department of Structural Biology, St. Jude Children’s Research Hospital, Memphis, TN, USA. ⁶Department of Cancer Biology, Dana-Farber Cancer Institute, Boston, MA, USA. ⁷Amazon Web Services, Seattle, WA, USA. ⁸Department of Pharmacy and Pharmaceutical Sciences, St. Jude Children’s Research Hospital, Memphis, TN, USA. ⁹Department of Biological Chemistry and Molecular Pharmacology, Harvard Medical School, Harvard University, Boston, MA, USA. ¹⁰Department of Biochemistry, University of Oxford, Oxford, UK. ¹¹Department of Bioengineering, Ege University, Bornova, Türkiye. ¹²Department of Chemistry, University of Zurich, Zurich, Switzerland. ¹³V. P. Kukhar Institute of Bioorganic Chemistry and Petrochemistry, National Academy of Science of Ukraine, Kyiv, Ukraine. ¹⁴Google, Mountain View, CA, USA. ¹⁵Department of Chemical Biology and Therapeutics, St. Jude Children’s Research Hospital, Memphis, TN, USA. ¹⁶Chemical Biology Program, Dana-Farber Cancer Institute, Boston, MA, USA. ¹⁷Zuse Institute Berlin (ZIB), Berlin, Germany. ¹⁸Max Delbrück Center for Molecular Medicine, Berlin, Germany. ¹⁹Department of Molecular Genetics, University of Toronto, Toronto, Ontario, Canada. ²⁰Vector Institute for Artificial Intelligence, Toronto, Ontario, Canada. ²¹Department of Chemistry, University of Toronto, Toronto, Ontario, Canada. ²²Department of Computer Science, University of Toronto, Toronto, Ontario, Canada. ²³Mathematical Institute, University of Oxford, Oxford, UK. ²⁴Department of Oncology, University of Oxford, Oxford, UK. ²⁵Harvard College, Cambridge, MA, USA. ²⁶Department of Biosciences and Bioengineering, Indian Institute of Technology Guwahati, Guwahati, India. ²⁷Department of Genetics, Stanford University, Stanford, CA, USA. ²⁸Department of Medical Oncology, Dana-Farber Cancer Institute, Boston, MA, USA. ²⁹Enamine Ltd, Kyiv, Ukraine. ³⁰Institute for Mathematics, Technical University Berlin, Berlin, Germany. ³¹NVIDIA, Santa Clara, CA, USA. ³²Intel, Santa Clara, CA, USA. ³³Wesleyan University, Middletown, CT, USA. ³⁴Pulse for Integrated Solutions, Unterschleißheim, Germany. ³⁵Canadian Institute for Advanced Research, Toronto, Ontario, Canada. ³⁶Department of Materials Science & Engineering, University of Toronto, Toronto, Ontario, Canada. ³⁷Institute of Medical Science, University of Toronto, Toronto, Ontario, Canada. ³⁸NVIDIA, Toronto, Ontario, Canada. ³⁹Acceleration Consortium, Acceleration Consortium, Toronto, Ontario, Canada. ⁴⁰Department of Chemical Engineering & Applied Chemistry, University of Toronto, Toronto, Ontario, Canada. ⁴¹Taras Shevchenko National University of Kyiv, Kyiv, Ukraine. ⁴²Chemspace, Kyiv, Ukraine. ⁴³Department of Physics, Faculty of Arts and Sciences, Harvard University, Cambridge, MA, USA. ⁴⁴Broad Institute of MIT and Harvard, Cambridge, MA, USA. ⁴⁵These authors contributed equally: Domiziana Cecchini, AkshatKumar Nigam, Ming Tang, Joana Reis.

✉ e-mail: andrea.mattevi@unipv.it; hari@hms.harvard.edu; christoph.gorgulla@stjude.org

Methods

AdaptiveFlow

AdaptiveFlow is an *in silico* platform consisting of three modules (AFLP, AFVS and AFU) and readymade ligand libraries for virtual screens. In this section, general platform-wide features are described that span multiple modules.

Python codebase. AdaptiveFlow has been entirely rewritten in Python, a programming language that enables the implementation of more sophisticated and advanced code while reducing maintenance work. This rewrite has also made it easier for scientists and programmers to join the open-source project and/or extend AdaptiveFlow independently, as Python is the most widely used programming language in scientific communities. In contrast, the previous version of AdaptiveFlow was written in Bash.

Parallelization. In addition to being rewritten in Python, a key goal for AdaptiveFlow was to improve parallelism to decrease the overall time required for screening. AdaptiveFlow has increased the levels of parallelism compared to the previous version. To help manage this increased parallelism, it maintains the use of collections that group around 1,000 ligands together, reducing the overhead of scheduling individual ligands for processing.

AdaptiveFlow also introduces the concept of subjobs, which are atomic units of computational work that include one or more collections for processing. These subjobs are automatically generated by AdaptiveFlow to match the runtime requirements specified in the job configuration. Subjobs are then combined into larger groups called work units, which are designed to allow schedulers to efficiently process the high levels of parallelism generated by AdaptiveFlow. These work units are sent as job arrays to the end schedulers for processing. By optimizing for parallelism, the team was able to achieve 5.6 million concurrent vCPUs running AdaptiveFlow.

Native cloud support. The AFLP and AFVS modules of AdaptiveFlow can run on computer clusters or on the cloud by using a batch system. Currently, two batch systems are supported, Slurm and AWS Batch. Slurm is able to run on many cloud infrastructures, including Google Cloud, Microsoft Azure and AWS. Within AWS, AdaptiveFlow uses AWS Batch and S3 object storage. CloudFormation is used to set up AdaptiveFlow and all necessary components. The AWS-native architecture allows AdaptiveFlow to use millions of vCPUs in parallel, while the previous version was only demonstrated to scale up to 160,000 vCPUs.

AdaptiveFlow is able to use various features from AWS in addition to supporting traditional clusters, making it easier to use and increasing the scale of computations available to researchers. Some of these features are as follows:

- **Scheduling and orchestration.** AWS Batch is a cloud-native scheduler and orchestrator for containerized jobs. As part of running a job, AdaptiveFlow can make use of AWS Batch to automatically scale cloud infrastructure up and down as needed and run all of the individual subjobs. AdaptiveFlow can also automatically generate the Docker container used for processing, eliminating the need to understand containerization or other infrastructure.
- **Object storage (Amazon S3).** In addition to the usage of standard cluster filesystems, AdaptiveFlow supports the usage of object storage for both ingesting of input datasets and the output of the resulting data from the processing.
- **Spot capacity awareness and spot resilience.** Spot capacity is unused computing capacity that cloud providers make available at greatly discounted prices compared to normal on-demand access. In exchange for a lower price, a workload running on a spot instance may be reclaimed (stopped) with a short warning.

AdaptiveFlow is designed to take advantage of this capacity model by using shorter running subjobs (approximately 30 min) and automatically rerunning jobs that have prematurely exited through integration with AWS Batch. Additionally, through integration with AWS Batch, AdaptiveFlow is able to select instances that are less likely to be interrupted through the 'spot-capacity-aware' scheduling method.

The architecture of AdaptiveFlow when deployed on AWS is shown in Supplementary Fig. 1. The architecture makes use of multiple compute environments, each of which may contain thousands of instances (compute nodes). The instances are automatically allocated and deallocated when the job starts and ends, respectively. Individual subjobs are allocated to the instances by AWS Batch until all subjobs have been completed. These subjobs run as Docker containers on the instances.

ARM. AdaptiveFlow is also able to run on the compute ARM architecture when using ARM-compatible choices in the settings. AFLP is fully ARM compatible and AFVS is ARM compatible when docking programs are used that either make the source code available or provide binaries for ARM. ARM-based computing infrastructure can be more cost-efficient than traditional x86-based architectures.

Enamine REAL Space

The REAL Space (enumerated, version 2022q1-2) was prepared using AFVS. JChem Suite was used for molecule neutralization, stereoisomer enumeration, tautomer generation, protonation state prediction, 3D coordinate calculation and the calculation of multiple molecular properties (JChem Suite 18.20.0; ChemAxon, <https://www.chemaxon.com>). Open Babel⁵⁶ (version 2.3.2) was used as a backup program for protonation state prediction, 3D coordinate calculation and target format conversion (all formats except SELFIES). The pip-installable SELFIES package (<https://github.com/aspuru-guzik-group/selfies>) was used for the conversion of molecules between SMILES and SELFIES. An overview of the processing steps and preparation methods used is shown in Supplementary Table 2.

Multiple formats. AdaptiveFlow provides the REAL Space library in ready-to-dock 3D formats including PDB, PDBQT, MOL2 and SDF. These formats are compatible with a wide range of docking programs for small molecules. The total size of the ready-to-dock library is approximately 50 TB in compressed form for each of these formats, totaling around 200 TB in compressed form (uncompressed, the library is around 2 PB). In addition to these ready-to-dock formats, the REAL Space library is also provided in the SELFIES³¹, SMILES⁵⁷ and Parquet formats.

AFLP

AFLP is the module of AdaptiveFlow dedicated to the preparation of ligand libraries into a ready-to-dock format and can handle libraries of any size, including ultralarge libraries. AFLP requires a batch system and supports Slurm and AWS Batch as batch systems. AFLP is configured using a configuration file (example shown in the Supplementary Information). Details on the workflow are available in the Supplementary Information. The new version of AFLP has the below-described features.

Preparation into ready-to-dock format. The primary task of AFLP is to prepare ligands into a ready-to-dock format. To accomplish this, AFLP desalts the original ligands (in case they are salts), neutralizes them, generates stereoisomers and tautomers, predicts the protonation states, computes the minimum-potential energy conformation of the ligand (3D coordinates), optionally validates the generated 3D conformers and converts them into the desired target formats. In total, over 100 output formats are supported. A detailed overview of the processing steps can be seen in Supplementary Fig. 3. The support input formats are SMILES, SELFIES and amino acid sequences. AFLP

uses external tools such as ChemAxon's JChem, Open Babel and RDKit (<https://www.rdkit.org/>) to carry out the preparation steps. For many of the preparation steps, the user can choose which external tool should be used. Details on which tools are available for each processing step can be found in Supplementary Table 3.

Molecular property calculation. AFLP enables researchers to calculate a variety of physicochemical properties and molecular descriptors for the molecules in the input ligand library, in addition to its primary task of preparing them for molecular docking. By calculating properties such as quantitative estimate of drug likeness, $\log P$, molecular weight and others (full list of supported properties in Supplementary Table 2), researchers can more effectively select and prioritize ligands, which can help reduce the cost and time required for screens using AFVS. AFLP uses mostly external tools to calculate the molecular properties, such as ChemAxon's JChem (<https://chemaxon.com/calculators-and-predictors>), Open Babel and RDKit (details in Supplementary Table 3).

Automatic tranche assignments. In AdaptiveFlow, ligand libraries are organized into tranches. Each tranche corresponds to a property of the ligands and is divided into multiple intervals or categories. For example, the molecular weight of the ligands can be used as a tranche type and then partitioned into several categories or intervals such as 0–200, 200–300, 300–400, 400–500 and >500 Da. Multiple tranches allow for the partitioning of a ligand library into a multidimensional grid or table, with one dimension corresponding to each tranche type (ligand property).

While the previous version of AdaptiveFlow required that the input ligand collection was already organized and structured in the tranche format, AdaptiveFlow can automatically reorganize the input ligand library into a new tranche output format on the basis of the computation of specific molecular properties of the ligands, as described above. This feature is useful because most ligand libraries (for example, from chemical compound vendors) are initially in the SMILES format and not in the tranche format. AFLP can then be used to prepare the entire library in the multidimensional grid or tranche format for AFVS.

To characterize how the Enamine REAL Space library populates the 18-dimensional property grid used by AdaptiveFlow, we quantified tranche occupancy (the number of molecules assigned to each grid cell) and summarized it as a histogram (Supplementary Fig. 37). We counted molecules per tranche across the full library. The resulting occupancy counts span several orders of magnitude; thus, we report the distribution using logarithmically spaced occupancy ranges (1, 2–10, 11–100, 101–1,000, 1,000–10,000, 10,000–100,000, 100,000–1 million, 1 million–10 million and >10 million molecules per tranche). The histogram displays the number of tranches in each occupancy range (y axis), highlighting that only a small fraction of the theoretical grid cells are populated and that populated tranches vary widely in size (from single-molecule tranches to highly populated tranches exceeding 10 million molecules), with an average of 5,600 molecules per populated tranche.

Support of SELFIES. SELFIES^{31,32} has become one of the most popular line notation formats for small molecules within ML communities. The primary reason is that SELFIES are fully robust, meaning that every SELFIES string corresponds to a valid molecule, which is not the case for SMILES. Thus, generative ML models cannot create invalid molecules when using SELFIES. AFLP now supports SELFIES as an input format for the molecules in addition to SMILES. Additionally, the entire REAL Space is free to be downloaded in the SELFIES format, enabling the creation of powerful datasets.

Stereoisomer enumeration. One of the main tasks of AFLP is to prepare ligand libraries from unprepared formats (for example, SMILES) into ready-to-dock formats. The new version of AFLP enables the

enumeration of stereoisomers for the ligands, which was not possible in the previous version. Accurate stereochemistry is essential for virtual screenings using structure-based molecular docking but many ligand libraries in unprepared formats (for example, from chemical compound vendors) do not contain stereochemical information. The ability to enumerate stereoisomers allows AFLP to include this information in the prepared ligand library, ensuring that the virtual screening results are accurate.

Extended open-source support. The previous version of AFLP required external tools (from ChemAxon) that are not open source for some of the steps when preparing ligands in a ready-to-dock format. The new version provides the option to choose open-source alternatives for each of the steps to prepare molecules into a ready-to-dock format. In particular, Open Babel and RDKit are supported by AdaptiveFlow as open-source alternatives.

Enhanced speed. AFLP has been improved in terms of speed compared to the original version. Firstly, this has been achieved by AFLP using its own Java code to call the necessary functions of ChemAxon's JChem package through Nailgun, which greatly reduces the number of calls to Nailgun. Additionally, timeout parameters that can be adjusted by users have been added to the control file, allowing AFLP to skip ligands that require unusually long processing times during preparation. These improvements make AFLP more efficient, robust and able to handle larger and more chemically complex ligand libraries.

Enhanced quality checks. AFLP checks the validity of each ligand that it prepares in the ready-to-dock format by calculating the potential energy of the ligand in 3D format using Open Babel. This allows AFLP to automatically remove ligands that have an excessively high potential energy, which can be an indication that the compound was damaged or corrupted during the preparation process. Additionally, AFLP supports optional plausibility checks on generated 3D conformers using PoseBusters⁵⁸. This helps to ensure the integrity and accuracy of the prepared ligand library.

Application. AFLP was used to prepare the REAL Space (version 2022q1-2). It was run on AWS using over 5.6 million Intel vCPUs in parallel. Spot instances were used and the used capacity was sustained without notable preemption. In total, <0.1% of the used CPU hours were interrupted because of preemption. The scaling behavior of AFLP was found to be perfectly linear (as shown in Supplementary Fig. 2). This demonstrates the efficiency and effectiveness of AFLP in preparing large ligand libraries in a ready-to-dock format.

AFVS

AFVS is the module of AdaptiveFlow dedicated to ULVSs but can also screen libraries of smaller sizes. It requires a batch system and supports Slurm and AWS Batch. AFVS is configured using a configuration file (example in the Supplementary Information). AFVS and AFU share many virtual screening features, which are described in a later section. AFVS and AFLP share the same parallelization mechanisms when using Slurm or AWS Batch; therefore, the scaling behavior of AFVS is expected to be the same as for AFLP, which was shown to be perfectly linear up to 5.6 million CPUs (Supplementary Fig. 2).

Parquet output format. AFVS has the ability to store output data in the Apache Parquet format, an open-source column-oriented data storage format. When using this format, AFVS can take advantage of any database and query system that can read this format, including Amazon Athena, a serverless service that allows users to efficiently query output score files stored in Amazon S3 (object storage) on AWS. Additionally, AFVS has a feature for automatically postprocessing ATG prescreens using Amazon Athena, enabling researchers to efficiently

prepare the input data for the ATG primary screens. Tools are provided with AFVS to facilitate common queries but researchers can also create their own queries using the open-source format of the dataset.

ATG-VSs and benchmarks. Following completion of the ATG prescreen, tranches can be selected for the ATG primary screen using two main strategies. In the first, individual tranches from the 18-dimensional REAL Space matrix can be manually selected, offering fine-grained control over specific chemical subspaces. In the second, selection can be restricted to a hyperrectangular region of the grid. This is achieved by identifying the top-performing intervals (that is, bins) within each of the 18 molecular property dimensions and computing their Cartesian product to form a structured, intersecting subset of tranches. This hyperrectangle approach enables chemically diverse yet targeted selection based on multiple favorable properties. In both selection modes, users define the total number of molecules to include in the ATG primary screen and AdaptiveFlow automatically prioritizes the tranches containing the best-performing representative ligands.

To evaluate the effectiveness of ATG-VS, we performed two types of benchmark studies: (1) full-scale production benchmarks using the entire Enamine REAL Space and (2) smaller-scale 'test-set-based' benchmarks designed for rapid prototyping and detailed comparison across different ATG configurations. In the test-set-based benchmarks, we randomly selected synthetic reactions from the Enamine space to create a compound subset totaling approximately 5 million molecules. For each of the ten protein targets tested (listed in Supplementary Table 7), a different random subset was used to ensure unbiased evaluation across targets and screening scenarios.

The test-set-based benchmarks allowed us to systematically assess the benefits of ATG-VS in three configurations: standard ULVS (random selection), ATG without active learning and ATG augmented with the optional ML-based active learning component. In contrast, the full-scale production benchmarks were conducted using the entire 69-billion-compound library, but without the active learning step, to minimize computational overhead while validating the performance of ATG-VS at scale. All docking in these benchmarks was performed using QuickVina 2 with default parameters, unless otherwise specified. For the 100K ATG benchmark configuration presented in Extended Data Fig. 5, tranches for the primary screen were selected using the hyperrectangle method, enabling a balanced, property-informed sampling of chemical space.

ML classifier for tranche prioritization. To improve the efficiency of ATG-VSs, the AdaptiveFlow platform incorporates an optional ML classification step into the primary screening stage. This step prioritizes compounds that are most likely to be high-affinity binders, thereby reducing the number of molecules subjected to full docking. The ML classifier is applied after the ATG prescreen and is used to filter undocked molecules within the selected tranches before docking.

The classifier is implemented as a fully connected feedforward neural network designed to distinguish between promising and less likely binders. Compounds are represented using Morgan fingerprints of length = 1,024 and radius = 2, generated by RDKit (version 2023.03.1). The model is trained on docking scores obtained from all representative molecules used in the prescreen. These scores are converted into binary classification labels using a percentile-based thresholding strategy: the top 25% of compounds with the most negative docking scores are labeled as positive examples ('high-confidence binders'), while the remaining compounds are labeled as negatives. To mitigate class imbalance, the majority class is undersampled during training.

The neural network architecture comprises four fully connected layers: an input layer of 1,024 nodes, followed by hidden layers of 512, 256 and 128 nodes, each using ReLU activation. The output layer consists of a single sigmoid-activated node that outputs the predicted binding probability. Hyperparameters and layer configurations were

selected through limited tuning using cross-validation on the prescreen data. The model is trained with the Adam optimizer (learning rate: 0.001), binary cross-entropy loss and early stopping based on validation loss with a patience of five epochs. A validation split of 10% is used, with training capped at 50 epochs and a batch size of 256, although convergence typically occurs earlier.

Once trained, the model is applied to all undocked compounds within the selected tranches. Compounds with predicted binding probabilities greater than 0.5 are retained for full docking, while the rest are filtered out. This filtering step typically reduces the number of compounds subjected to docking by 30–70%, depending on the target, while maintaining or improving hit enrichment.

The ML classification module is fully integrated into the AFU environment and can be activated through a configuration parameter. The training and inference pipeline is implemented in PyTorch (version 2.0.1) and supports GPU acceleration for scalable and efficient deployment.

AFU

AFU is the all-in-one version of AdaptiveFlow, combining both ligand preparation and virtual screening into a single streamlined workflow. It is the third module alongside AFLP and AFVS. This new module adds many conveniences for workflows where both ligand preparation and virtual screenings have to be carried out in tandem. Notably, AFU is designed to be run as a package, supporting approximately 1,500 docking protocols that can be run independently of job schedulers (Supplementary Table 8). A conceptual overview of the workflow of AFU can be found in Supplementary Fig. 8 and a detailed UML workflow diagram is shown in Supplementary Fig. 14. Specifically, as input, a user specifies the molecule to be docked (as a SMILES, SELFIES or amino acid sequence), the protein file, the docking site and the docking program to be run. On the basis of the user's choices, error checks in the form of missing docking-choice-specific files are performed and calculations are run. As a result of a successful calculation, the docking score and pose are returned. AFU can be run in standalone mode, where it is run through the command line and specifies the options using a configuration file. This mode is for example useful for scientists in the drug discovery communities who want to carry out molecular dockings in the most convenient and simple way, without requiring computational expertise, while having a large number of docking protocols available. In addition, AFU can be run in API mode to directly interface with other programs and code, which can, for instance, be useful for the ML communities that develop new methods for drug discovery.

AFU comes in two versions. The first is dedicated to single workstations or compute nodes without multinode parallelization. The second, AFUparr, is the parallelized version of VirtualFlow Unity. Designed to operate seamlessly on SLURM systems, AFUparr allows users to easily parallelize the workflow on a larger number of CPUs and GPUs. Execution within SLURM environments is highly optimized, with computations distributed in parallel across multiple CPUs and nodes. This design ensures efficient linear scaling relative to the number of molecules provided.

Joint features of AFVS and AFU

The AFU and AFVS modules support the same docking protocols and, therefore, share all docking-specific features.

Supported docking protocols. AdaptiveFlow supports approximately 40 docking protocols, listed in Supplementary Tables 4–6. These include software that performs both pose prediction and scoring (Supplementary Table 4), solely pose prediction (Supplementary Table 6) and only scoring (Supplementary Table 5). Notably, within AdaptiveFlow, the scoring functions can be combined with pose prediction methods, giving rise to approximately 1,500 methods (listed in Supplementary Table 8). We note that many of the newly supported

docking software possess special features such as protein–protein or RNA docking, explicit parameterization of solvents or metals and covalent docking. In Supplementary Tables 4–6, we describe these as ‘special features’. Additionally, a key consideration for ULVS is the timing required for carrying out a single calculation. As such, in Extended Data Fig. 9, we provide timing estimates for running a single calculation, averaging over 25 independent runs for all docking programs.

Special types of ligands. In addition to traditional small molecule ligands, AdaptiveFlow supports the use of peptides, protein ligands and covalently bonded ligands in AFVS. Peptides are of increasing interest as therapeutics and several docking programs already support their use. For protein–protein or RNA docking, specialized programs require two protein or RNA targets to generate co-complexes. Covalent docking involves reduced flexibility for certain parts of the input molecule during pose sampling and a few specialized docking programs support this ligand type (Supplementary Table 4).

Enhanced performance. The QuickVina 2 molecular docking tool is used to accelerate AutoDock Vina with no loss of accuracy²⁷. QuickVina 2 execution was profiled in an effort to improve performance even further. To establish baseline performance, the open-source GitHub repository was cloned and the `qvina_1buffer` branch was compiled using the GNU C++ compiler (version 9.4), the Boost C++ libraries (version 1.79, <https://www.boost.org>) and the default parameters (build instructions: <https://qvina.github.io/compilingQvina2.html>). A single docking experiment of avanafil binding to unsolvated PDB 5FNQ was run through the Intel VTune Profiler to look for performance bottlenecks. Unfortunately, the execution profile was flat. The most time-consuming function accounted for less than 20% of the total runtime. Tuning the source code would give rapidly diminishing returns at the risk of affecting accuracy; hence, compiler-level optimization was the best option to improve performance. QuickVina 2 was rebuilt using the Intel oneAPI DPC++/C++ Compiler (version 2022.2.0.20220730) and aggressive optimization: automatic vectorization, interprocedural optimization and relaxed floating-point constraints. Profile-guided optimization was also tried but it did not improve performance. The final executable was 18% faster than the default builds with no loss of accuracy. The new binaries are provided on the QuickVina and AF homepages.

GPU support. Several of the docking programs support GPUs, which can strongly reduce computation time and cost. For pose prediction, these include AutoDock-GPU⁵², QVina2-GPU⁵³, QVina2-W-GPU⁵³ and Vina=GPU⁵³. Gnina⁵⁹, NNScore (version 2.0)⁶⁰, DeepBindRG⁶¹ and DeepAffinity⁶² use GPU-compatible deep learning models for scoring protein–ligand complexes. Additionally, deep learning methods such as TANKBind¹², DiffDock¹³ and EquiBind²⁸ use GPU-compatible neural networks for direct pose prediction of several ligands in parallel. AdaptiveFlow supports GPU compatibility of all these software. The relative speedups of these methods for single calculations can be compared within Extended Data Fig. 9.

Molecular-dynamics-based methods. The combination of molecular docking and molecular dynamics can yield results near experimentally determined crystal structures^{63,64}. AFVS includes methods for selecting and scoring poses on the basis of molecular dynamics simulations, such as MM/PBSA and MM/GBSA⁶⁵, for evaluating docking poses and predicting binding affinities. MM/PBSA and MM/GBSA are popular for predicting binding free energy because they are more accurate than molecular docking and less computationally intensive than alchemical methods. AdaptiveFlow also supports binding pose metadynamics (BPMD), which uses a stability score based on the resistance of the ligand to perturbation away from the initial pose to rerank predicted binding poses from docking. The OpenBPMD algorithm⁶⁶ can accurately predict

binding poses within 2 Å (root-mean-square deviation) of the crystallographic complex structure for many systems and efficiently rerank docked poses. The MM/PBSA, MM/GBSA and BPMD methods can be used within AdaptiveFlow as rescoring and pose refinement methods.

Enhanced quality checks. AFVS supports optional plausibility checks on generated docking poses of Vina-based docking protocols using PoseBusters, an open-source tool, which checks for chemical as well as intramolecular ligand and intermolecular validity of ligand–protein complexes⁵⁸.

Reporting summary

Further information on research design is available in the Nature Portfolio Reporting Summary linked to this article.

Data availability

Additional data can be found on the AdaptiveFlow homepage (<http://adaptive-flow.ai/>). FSP1 structures were deposited to the Protein Data Bank under accession codes PDB 9IFZ, 9IFT, 9IFY, 9IFU, 9IFW, 9IG9 and 9QRT. The structure of the PARP1–inhibitor complex was deposited to the Protein Data Bank under accession code PDB 8VYH. The FSP1 ancestral sequence generated in this study was submitted for deposition to the GenBank database with accession code PQ619396. The ready-to-dock Enamine REAL Space can be explored online through an interactive web interface (<https://enamine.net/compound-collections/real-compounds/real-space-navigator>). The data can be accessed directly through an AWS S3 bucket provided by AWS Open Data^{67–71} (<https://aws.amazon.com/opendata>). Additional information is available on the AWS Open Data website of this dataset (https://aws.amazon.com/marketplace/pp/prodview-m32p7engdxk7k?sr=0-1&ref_=beagle&applicationId=AWSMPContessa), as well as on the AdaptiveFlow homepage (<https://adaptive-flow.ai/real-space-2022q12>; free registration required). The original (raw) SMILES of the REAL Space from Enamine are made available under the CC-BY attributed to Enamine. The data supporting the findings of this study are available in the paper and its Supplementary Information. Source data are provided with this paper.

Code availability

The source code of AdaptiveFlow is released under the GNU GPL (version 2.0), a free and open-source license. The code can be accessed through GitHub (<https://github.com/QuantumAI4Bio>), which contains multiple repositories. The AFLP and AFVS modules of AdaptiveFlow can be found in the AFLP (<https://github.com/QuantumAI4Bio/AdaptiveFlow-LP>) and AFVS (<https://github.com/QuantumAI4Bio/AdaptiveFlow-VS>) GitHub repositories, respectively. The AFU module (both single machine and parallelized versions) can be found in the GitHub repository (<https://github.com/QuantumAI4Bio/AdaptiveFlow-Unity>)^{67–71}.

References

- Buttenschoen, M., Morris, G. M. & Deane, C. M. PoseBusters: AI-based docking methods fail to generate physically valid poses or generalise to novel sequences. *Chem. Sci.* **15**, 3130–3139 (2024).
- McNutt, A. T. et al. GNINA 1.0: molecular docking with deep learning. *J. Cheminform.* **13**, 1–20 (2021).
- Durrant, J. D. & McCammon, J. A. NNScore 2.0: a neural-network receptor–ligand scoring function. *J. Chem. Inf. Model.* **51**, 2897–2903 (2011).
- Zhang, H., Liao, L., Saravanan, K. M., Yin, P. & Wei, Y. DeepBindRG: a deep learning based method for estimating effective protein–ligand affinity. *PeerJ* **7**, e7362 (2019).
- Karimi, M., Wu, D., Wang, Z. & Shen, Y. DeepAffinity: interpretable deep learning of compound–protein affinity through unified recurrent and convolutional neural networks. *Bioinformatics* **35**, 3329–3338 (2019).

63. Clark, A. J. et al. Prediction of protein–ligand binding poses via a combination of induced fit docking and metadynamics simulations. *J. Chem. Theory Comput.* **12**, 2990–2998 (2016).
64. Miller, E. B. et al. Reliable and accurate solution to the induced fit docking problem for protein–ligand binding. *J. Chem. Theory Comput.* **17**, 2630–2639 (2021).
65. Wang, E. et al. End-point binding free energy calculation with mm/pbsa and mm/gbsa: strategies and applications in drug design. *Chem. Rev.* **119**, 9478–9508 (2019).
66. Lukauskis, D. et al. Open binding pose metadynamics: an effective approach for the ranking of protein–ligand binding poses. *J. Chem. Inf. Model.* **62**, 6209–6216 (2022).
67. Gorgulla, C. et al. Ready-to-dock Enamine REAL Space (version 2022q1-2). *AWS Open Data Registry* <https://aws.amazon.com/marketplace/pp/prodview-m32p7engdxk7k> (2025).
68. Gorgulla, C. et al. AdaptiveFlow-LP: library preparation module v1.0.0. *Zenodo* <https://doi.org/10.5281/zenodo.20412411> (2026).
69. Gorgulla, C. et al. AdaptiveFlow-VS: virtual screening module v1.0.0. *Zenodo* <https://doi.org/10.5281/zenodo.20412363> (2026).
70. Nigam, A. K. et al. AdaptiveFlow-Unity (single-machine) v1.0.0. *Zenodo* <https://doi.org/10.5281/zenodo.20412397> (2026).
71. Nigam, A. K. et al. AdaptiveFlow-Unity (parallelized) v1.0.0. *Zenodo* <https://doi.org/10.5281/zenodo.20412415> (2026).

Acknowledgements

We would like to thank ChemAxon for a free academic license for JChem. A.K.N. thanks A. Kundaje and M. Bassik for their support. We thank AWS for research credits for the AWS Cloud and N. Gonzales, S. Litster, M. Thakkar, R. Wang, A. Subramaniyan and C. Tsien Silvers for technical and administrative support within AWS. We would like to thank A. Bilani and C. Rhoades from Intel for their support, B. Klein for proofreading the manuscript and E. Clifton regarding the formatting of the manuscript. We acknowledge the use of AI-based language tools (for example, ChatGPT, Claude and Gemini) for editing and language polishing during manuscript preparation.

Author contributions

C.G. conceptualized and led the development of the AdaptiveFlow computational platform, the ATG-VS method and the 18-dimensional chemical library, built and led the multi-institutional AdaptiveFlow development team, supervised the computational work and the virtual screens applied to FSP1 and PARP1, established the collaboration with AWS, initiated and coordinated the integration of the FSP1 experimental program into the AdaptiveFlow project and led the writing of the computational parts of the manuscript. A.K.N., M.T. and M.K. performed the dockings, screenings and benchmarks. A.J., S.S.C., Y.M., M.G., J.K., N.T., C.S., M.H., A.H., M.L., Abhishek Kumar, R.L.-K., Aditya Kumar, P.-Y.A., H.G., A.A., E.O., A.A.G. and K.F. contributed to computational method development and implementation. A.M. conceptualized and supervised the FSP1 experimental program and led the FSP1 sections of the manuscript. A.M., D.C. and A.G. performed all biochemical and structural analyses of FSP1. C.R.N. performed phylogenetic analysis and ancestral sequence reconstruction. J.D.S. supervised the cellular studies on FSP1, performed by Y.W. Y.L. and L.Y. performed the cellular ADME studies on FSP1. K.M.P.D. and K.C. performed the surface plasmon resonance experiments on FSP1. H.A. contributed to early design discussions and provided input during the development of AdaptiveFlow. H.A. led the experimental PARP1 ligand discovery program, directed compound selection from the secondary computational screen, designed the experimental validation pipeline

and supervised all PARP1 experimental work, including biochemical characterization, NMR validation, X-ray crystallography and cellular validation in BRCA1-deficient triple-negative breast cancer models, and led the PARP1 sections of the manuscript. J.R., R.T., H.-S.S., S.-D.P., H.C., L.S., C.B. and P.B. performed the biochemical and structural analyses of PARP1. E.-B.C. and G.I.S. performed the cell-based studies involving PARP1. D.S.R. and Y.S.M. contributed to the chemical syntheses. C.G.K. and G.W. provided supervision and contributed to scientific discussions. All authors provided critical feedback and helped to shape the research, analysis and manuscript.

Funding

A.K.N. acknowledges funding from the Bio-X Stanford Interdisciplinary Graduate Fellowship. C.S. and K.F. acknowledge funding by the Deutsche Forschungsgemeinschaft (German Research Foundation) under Germany's Excellence Strategy—The Berlin Mathematics Research Center MATH+ (EXC-2046/1, project ID: 390685689). H.A. gratefully acknowledges support from the National Institutes of Health National Institute of General Medical Sciences (GM136859 and GM158220), a grant from Pivotal Life Sciences and a gift from J. Goldberg through Dana-Farber. A.M. acknowledges funding from the European Research Council (Advanced Grant, MetaQ no. 101094471) and the Associazione Italiana per la Ricerca sul Cancro Investigator Grant (no. 28754). A.G. is supported by a fellowship from the Associazione Italiana per la Ricerca sul Cancro (no. 32882 in memory of Angelo and Giancarla Giardina). C.G. acknowledges funding from the CADET grant from St. Jude Children's Research Hospital and the Blue Sky Kinase Project at St. Jude Children's Research Hospital, as well as support from American Lebanese Syrian Associated Charities, the fundraising and awareness organization of St. Jude Children's Research Hospital.

Competing interests

D.S.R. works for Enamine, a company that is involved in the synthesis and distribution of drug-like compounds. Y.S.M. works for Onepot AI. C.G. and G.W. are cofounders of Virtual Discovery, which is a fee-for-service company for computational drug discovery, and C.G. also consults for Virtual Discovery. M.G. works for Google, a company providing cloud computing services. M.K. and P.-Y.A. work for AWS, a company providing cloud computing services. C.S. is the founder and owner of Phind Beteiligungs UG, a company providing consulting services in virtual screening and computational drug discovery. The remaining authors declare no competing interests.

Additional information

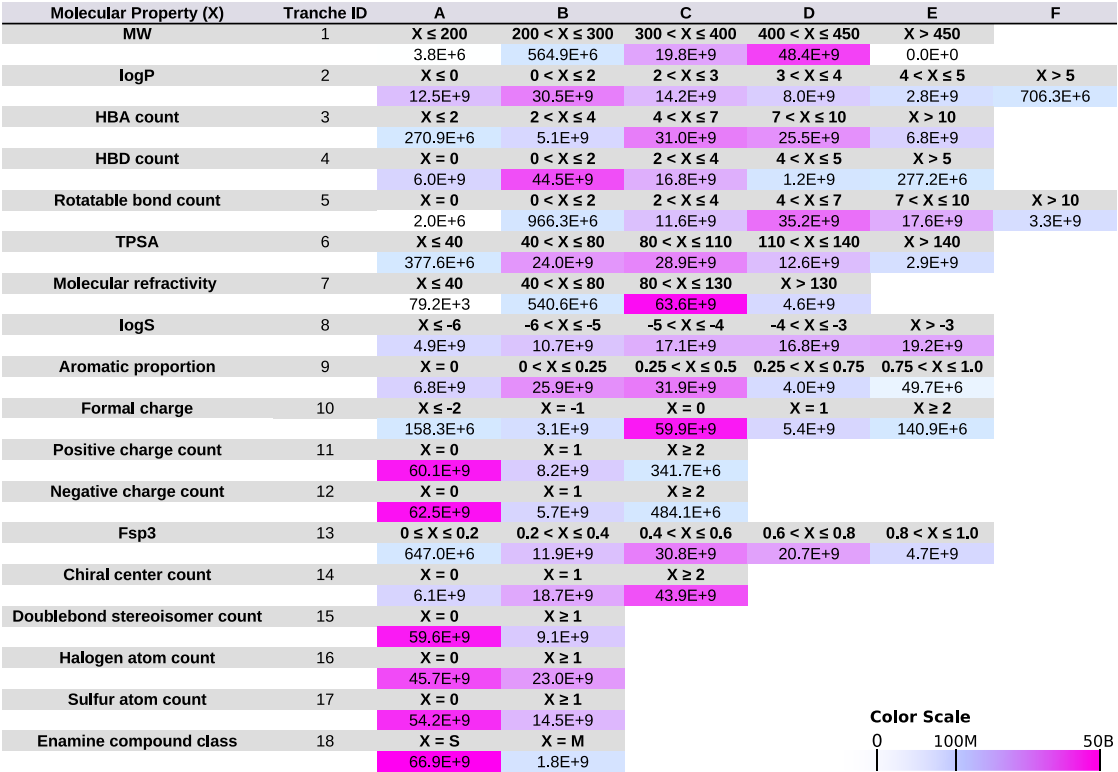
Extended data is available for this paper at <https://doi.org/10.1038/s41587-026-03217-x>.

Supplementary information The online version contains supplementary material available at <https://doi.org/10.1038/s41587-026-03217-x>.

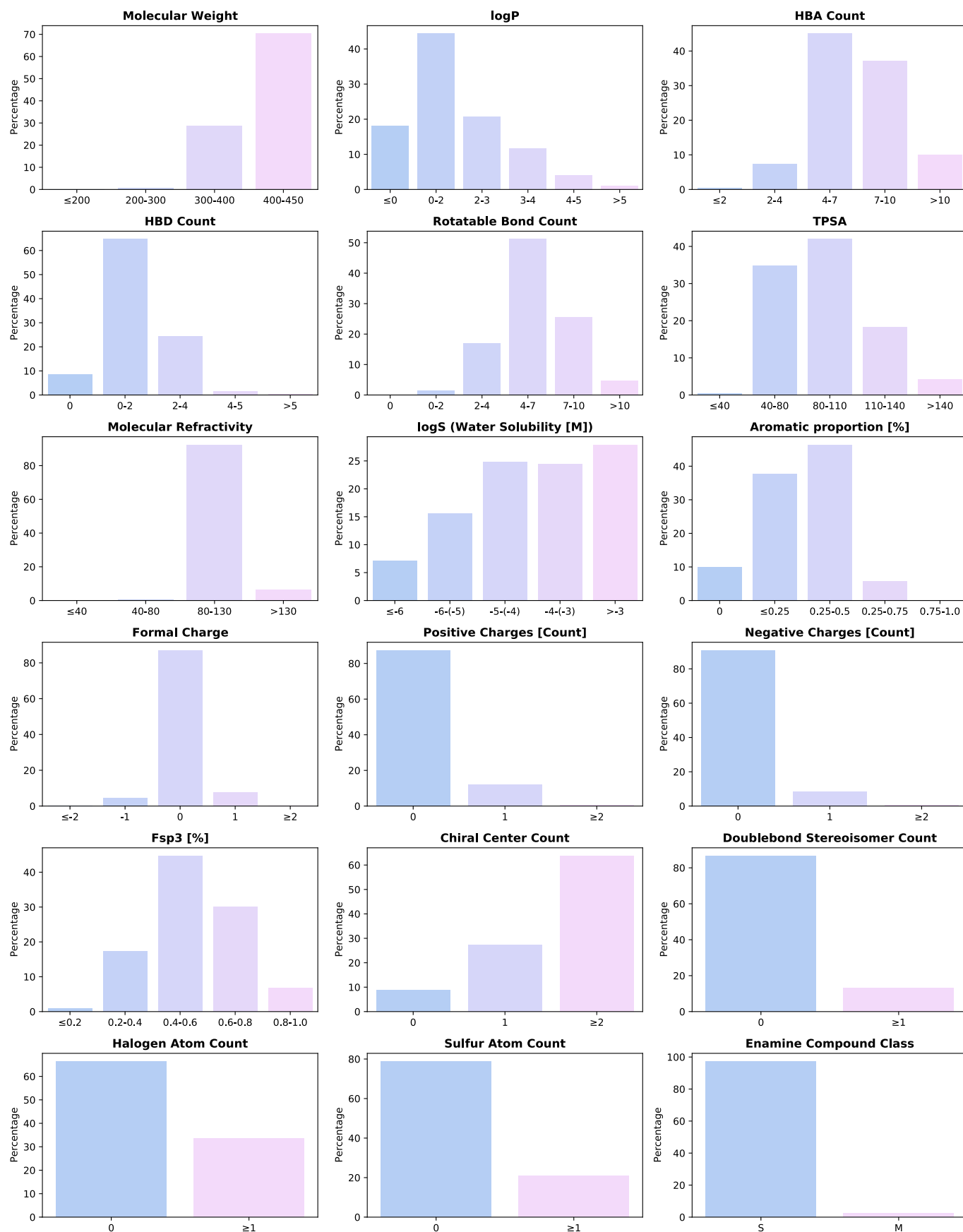
Correspondence and requests for materials should be addressed to Andrea Mattevi, Haribabu Arthanari or Christoph Gorgulla.

Peer review information *Nature Biotechnology* thanks Peter Kolb and the other, anonymous, reviewer(s) for their contribution to the peer review of this work.

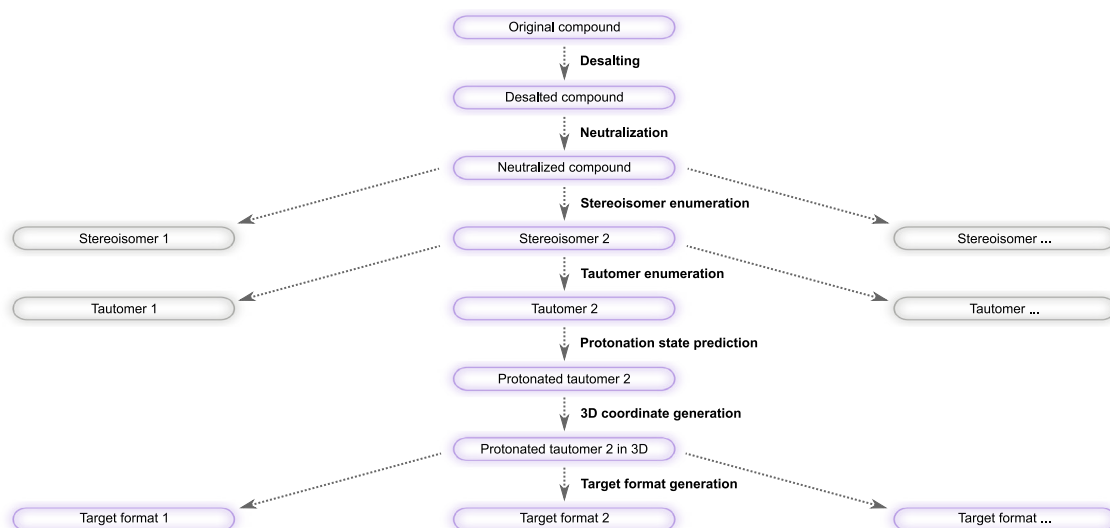
Reprints and permissions information is available at www.nature.com/reprints.



Extended Data Fig. 1 | Overview of 18-dimensional tranche table. Molecular properties associated with each tranche, the partitions that are used for each tranche property, and the number of compounds per partition interval for each property (colored as a heat map).



Extended Data Fig. 2 | Overview of the distribution of the molecular properties of the Enamine REAL Space (version 2022q1-2). Each subplot shows the distribution of the molecular property associated with each dimension of the 18-dimensional grid into which the REAL Space was partitioned.

**Extended Data Fig. 3 | AFLP - Overview of the Ligand Preparation Procedure.**

Ligands are initially in a format in line notation (e.g. SMILES, SELFIES, or amino acid sequences). The preparation steps include desalting, neutralization, stereoisomer enumeration, tautomer enumeration, protonation state prediction, 3D coordinate generation, and target format generation.

Stereoisomer and tautomer generation can result in multiple chemical species. Highlighted in purple are the steps of a molecule from the beginning to the end of the preparation procedure. Colored in white are the intermediate molecules/tautomers whose subsequent preparation steps are not shown in the diagram.

REAL Space (2022q12) – Interactive Tranche Table

View Edit Outline Delete Revisions

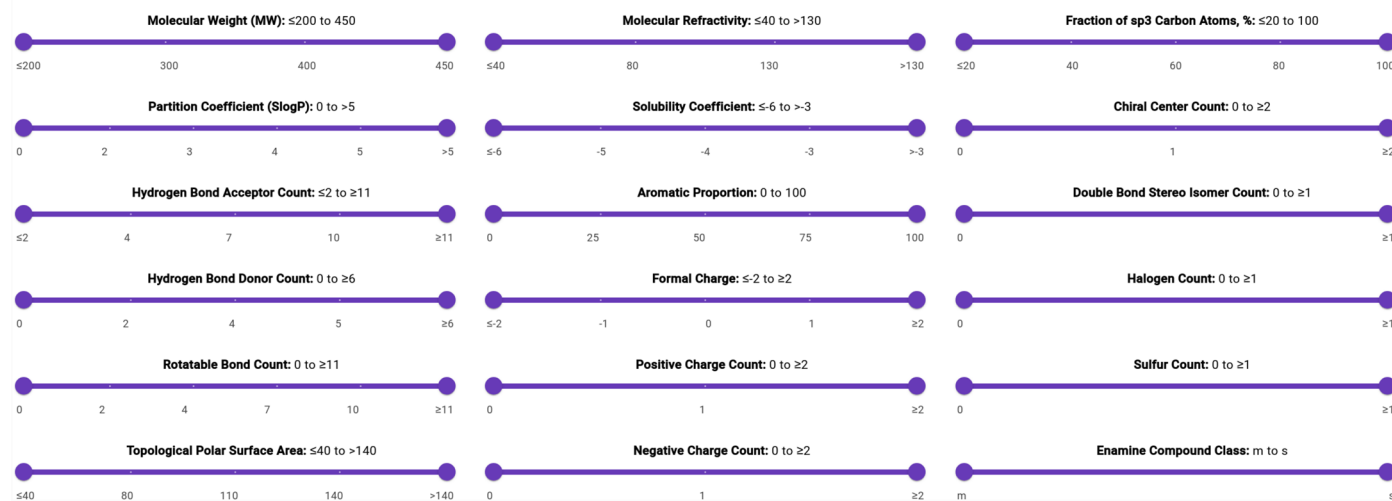
Vertical dimension
Molecular Weight (MW)

	≤200	300	400	450	Row sums:
0	1M	120.8M	4.4B	8B	12.5B
2	2M	264.6M	9.3B	20.9B	30.5B
3	512.9K	103.6M	3.6B	10.5B	14.2B
4	157.4K	53.6M	1.7B	6.2B	8B
5	17.2K	18.4M	552.5M	2.3B	2.8B
>5	253	4M	149.2M	553.1M	706.3M
Column sums:	3.8M	564.9M	19.8B	48.4B	

Horizontal dimension
Partition Coefficient (SlogP)

Compounds selected: 68.7B from 68.7B Tranches selected: 12.3M from 12.3M

Get script



Extended Data Fig. 4 | Webinterface (version 2) of the REAL Space. Shown is an interactive web interface that is made available on the AdaptiveFlow homepage (<https://adaptive-flow.ai/real-space-2022q12-interactive-tranche-table>) that allows selecting subsets of the REAL Space partitioned into an 18-dimensional matrix. Each dimension corresponds to a molecular property and has a slider that

can be adjusted to select the desired property ranges. The number of selected molecules is updated automatically in the web interface. Effectively, the sliders allow the selection of multidimensional rectangles within the 18-dimensional matrix.

Property (X)	Tranche		Reps/Tranche=1, Exh=1		Reps/Tranche=10, Exh=1		Reps/Tranche=1, Exh=10	
	Letter	Range	Average Score	Rank	Average Score	Rank	Average Score	Rank
Molecular weight	A	$X \leq 200$	-5.69868	4	-5.66977	4	-5.75664	4
	B	$200 < X \leq 300$	-7.2417	3	-7.20452	3	-7.31959	3
	C	$300 < X \leq 400$	-7.79181	1	-7.76165	1	-7.95967	2
logP	D	$400 < X \leq 450$	-7.76098	2	-7.74938	2	-8.0356	1
	A	$X \leq 0$	-7.67313	5	-7.66211	6	-7.86196	6
	B	$0 < X \leq 2$	-7.68707	4	-7.6674	5	-7.87787	5
	C	$2 < X \leq 3$	-7.7057	3	-7.68094	3	-7.90026	4
	D	$3 < X \leq 4$	-7.71324	2	-7.69173	2	-7.90955	2
	E	$4 < X \leq 5$	-7.7178	1	-7.70489	1	-7.92332	1
	F	$X > 5$	-7.67281	6	-7.66768	4	-7.90419	3
HBA count	A	$X \leq 2$	-7.22028	5	-7.20911	5	-7.32189	5
	B	$2 < X \leq 4$	-7.38033	4	-7.36258	4	-7.53528	4
	C	$4 < X \leq 7$	-7.61231	3	-7.59707	3	-7.8139	3
	D	$7 < X \leq 10$	-7.79858	2	-7.78421	2	-8.01877	2
HBD count	E	$X > 10$	-8.00413	1	-7.9967	1	-8.22418	1
	A	$X = 0$	-7.6844	3	-7.66601	4	-7.86789	4
	B	$0 < X \leq 2$	-7.70107	2	-7.68555	1	-7.89879	2
	C	$2 < X \leq 4$	-7.70723	1	-7.68371	2	-7.90943	1
	D	$4 < X \leq 5$	-7.68301	4	-7.6664	3	-7.88586	3
Rotatable bond count	E	$X > 5$	-7.63874	5	-7.62658	5	-7.83489	5
	A	$X = 0$	-7.71737	4	-7.7189	4	-7.85219	4
	B	$0 < X \leq 2$	-7.94585	2	-7.90628	2	-8.06441	2
	C	$2 < X \leq 4$	-7.99004	1	-7.9602	1	-8.11996	1
	D	$4 < X \leq 7$	-7.81481	3	-7.79696	3	-7.97132	3
TPSA	E	$7 < X \leq 10$	-7.35753	5	-7.36396	5	-7.6019	5
	F	$X > 10$	-6.72586	6	-6.74264	6	-7.2096	6
	A	$X \leq 40$	-7.18856	5	-7.18394	5	-7.2971	5
	B	$40 < X \leq 80$	-7.489	4	-7.47825	4	-7.66254	4
	C	$80 < X \leq 110$	-7.72677	3	-7.71709	3	-7.94197	3
Molecular refractivity	D	$110 < X \leq 140$	-7.90033	2	-7.88201	2	-8.13229	2
	E	$X > 140$	-8.10076	1	-8.08286	1	-8.32545	1
	A	$X \leq 40$	-5.44523	4	-5.38075	4	-5.52602	4
	B	$40 < X \leq 80$	-7.31169	3	-7.25362	3	-7.40034	3
	C	$80 < X \leq 130$	-7.73766	2	-7.71934	2	-7.9346	2
logS	D	$X > 130$	-7.79962	1	-7.79242	1	-8.07744	1
	A	$X \leq -6$	-7.89126	1	-7.87113	1	-8.125	1
	B	$-6 < X \leq -5$	-7.83936	2	-7.82017	2	-8.04633	2
	C	$-5 < X \leq -4$	-7.75784	3	-7.73812	3	-7.95445	3
	D	$-4 < X \leq -3$	-7.67255	4	-7.65238	4	-7.86527	4
Aromatic proportion	E	$X > -3$	-7.56383	5	-7.55104	5	-7.75079	5
	A	$X = 0$	-7.00124	5	-6.97872	5	-7.2258	5
	B	$0 < X \leq 0.25$	-7.47794	4	-7.4629	4	-7.71839	4
	C	$0.25 < X \leq 0.5$	-7.69386	3	-7.687	3	-7.88349	3
	D	$0.25 < X \leq 0.75$	-8.07776	2	-8.07503	2	-8.23953	2
Formal charge	E	$0.75 < X \leq 1.0$	-8.38983	1	-8.39619	1	-8.51087	1
	A	$X \leq -2$	-8.26463	1	-8.28754	1	-8.43463	1
	B	$X = -1$	-7.95327	2	-7.94973	2	-8.1436	2
	C	$X = 0$	-7.66391	3	-7.64593	3	-7.85809	3
	D	$X = 1$	-7.53418	4	-7.50755	4	-7.7401	4
Positive charge count	E	$X \geq 2$	-7.26159	5	-7.21597	5	-7.4746	5
	A	$X = 0$	-7.64767	3	-7.62676	3	-7.83845	3
	B	$X = 1$	-7.73326	2	-7.72519	2	-7.93725	2
Negative charge count	C	$X \geq 2$	-7.7854	1	-7.78236	1	-7.9847	1
	A	$X = 0$	-7.47956	3	-7.46405	3	-7.677	3
	B	$X = 1$	-7.85469	2	-7.86841	2	-8.05519	2
Fsp3	C	$X \geq 2$	-8.23447	1	-8.26297	1	-8.41197	1
	A	$0 \leq X \leq 0.2$	-8.25291	1	-8.24892	1	-8.39246	1
	B	$0.2 < X \leq 0.4$	-7.9144	2	-7.91246	2	-8.09817	2
	C	$0.4 < X \leq 0.6$	-7.57134	3	-7.5644	3	-7.78127	3
	D	$0.6 < X \leq 0.8$	-7.27781	4	-7.26214	4	-7.50539	4
Chiral center count	E	$0.8 < X \leq 1.0$	-7.02762	5	-6.99253	5	-7.26721	5
	A	$X = 0$	-7.68648	2	-7.66472	3	-7.88864	2
	B	$X = 1$	-7.68156	3	-7.66601	2	-7.87407	3
Doublebond stereoisomer count	C	$X \geq 2$	-7.72642	1	-7.70654	1	-7.91929	1
	A	$X = 0$	-7.7409	1	-7.72163	1	-7.92645	1
Halogen atom count	B	$X \geq 1$	-7.58906	2	-7.56916	2	-7.81019	2
	A	$X = 0$	-7.6665	2	-7.65526	2	-7.88305	2
Sulfur atom count	B	$X \geq 1$	-7.7349	1	-7.70825	1	-7.90392	1
	A	$X = 0$	-7.79078	1	-7.77764	1	-7.98746	1
Enamine compound class	B	$X \geq 1$	-7.5662	2	-7.53497	2	-7.7627	2
	A	$X = S$	-7.67793	2	-7.65586	2	-7.89126	2
	B	$X = M$	-7.72188	1	-7.71391	1	-7.89299	1

Extended Data Fig. 5 | See next page for caption.

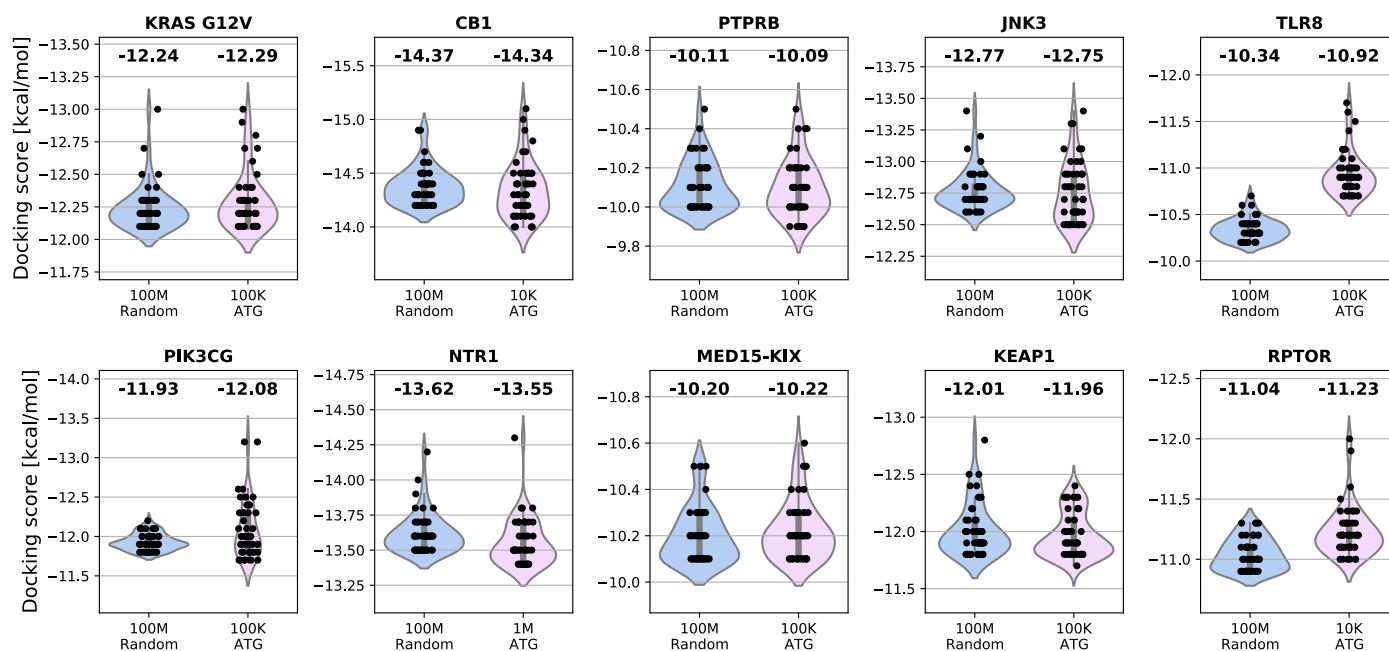
Extended Data Fig. 5 | Adaptive Target-Guided Prescreening Benchmarks.

Results of several prescreens for a test system within the Adaptive Target-Guided Virtual Screening (ATG-VS) approach, where different settings were used.

In the first run (left), one representative per tranche was used with docking exhaustiveness 1. In the second run (middle), 10 representatives per tranche are used with exhaustiveness 1. In the third run (right), one representative per

tranche was used with exhaustiveness 10. The average docking scores were computed for each tranche of each property (averaging over all other properties/dimensions). For each property, the tranches have been ranked within their property realm based on the average docking scores for the tranches. As can be seen, the patterns and rankings are very similar to each other, indicating that one representative with exhaustiveness 1 can be sufficient for the prescreenings.

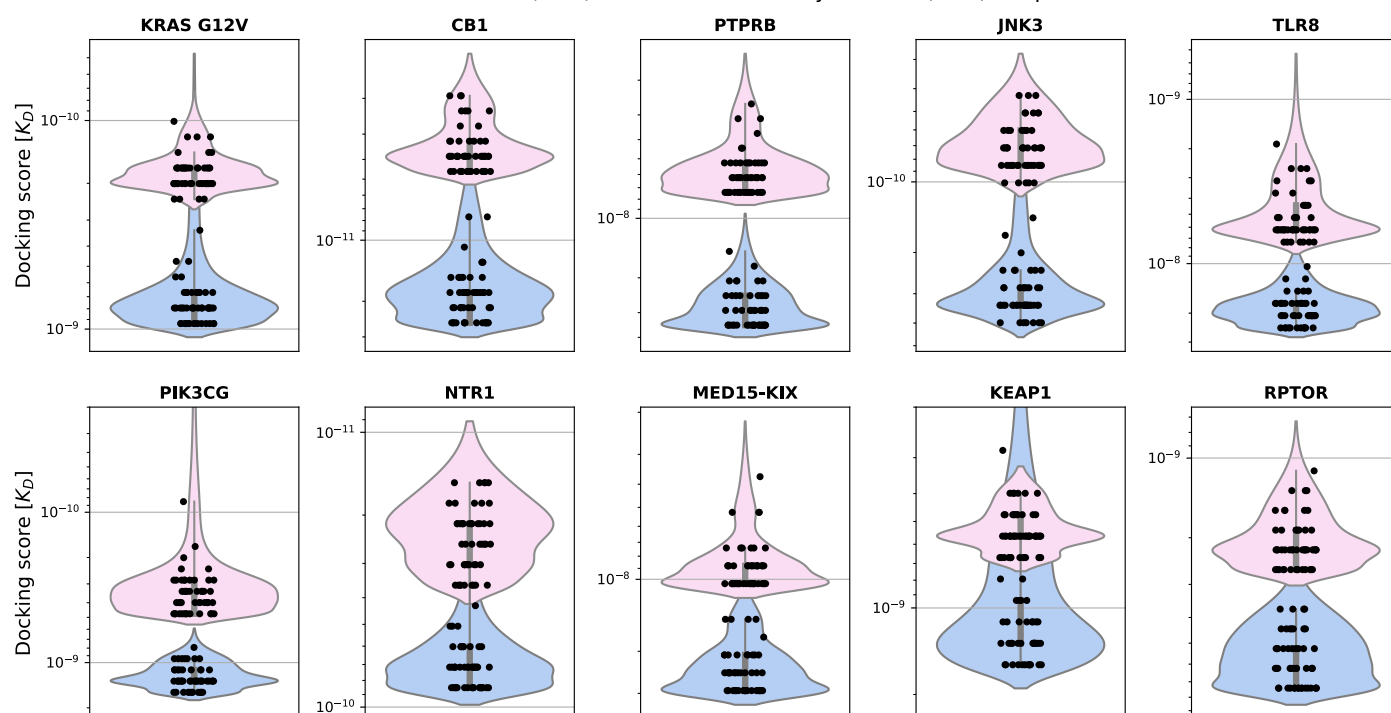
100M Random (blue) vs Adaptive Target Guided (ATG) Primary Screens - Top 50 compounds

**Extended Data Fig. 6 | Benchmarking large-scale virtual screening strategies.**

Violin plots compare the distribution of docking scores for the top 50 molecules identified by two approaches¹: a standard ultra-large virtual screen of 100 million randomly selected compounds, and² the Adaptive Target-Guided (ATG) screening strategy developed in this work. For each ATG run, a prescreen of 11

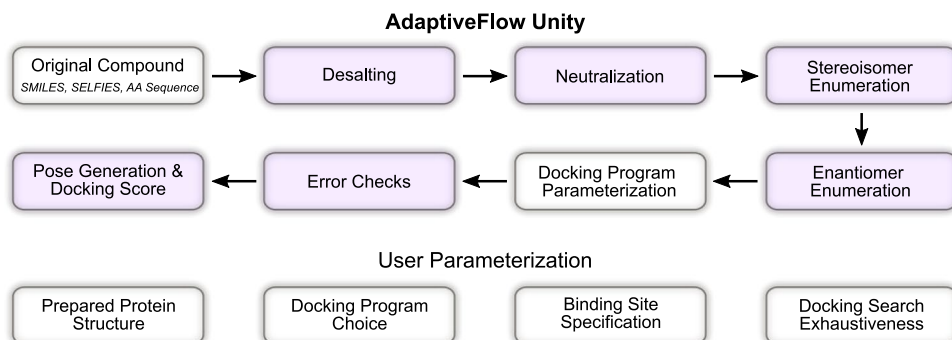
million representative compounds was used to identify the most promising regions of chemical space, followed by a focused primary screen whose size is indicated for each target. ATG consistently achieves comparable or superior docking scores using a fraction of the computational budget.

100M Standard ULVSs (blue) vs 100M ATG Primary Screens (rose) - Top 50 Hits



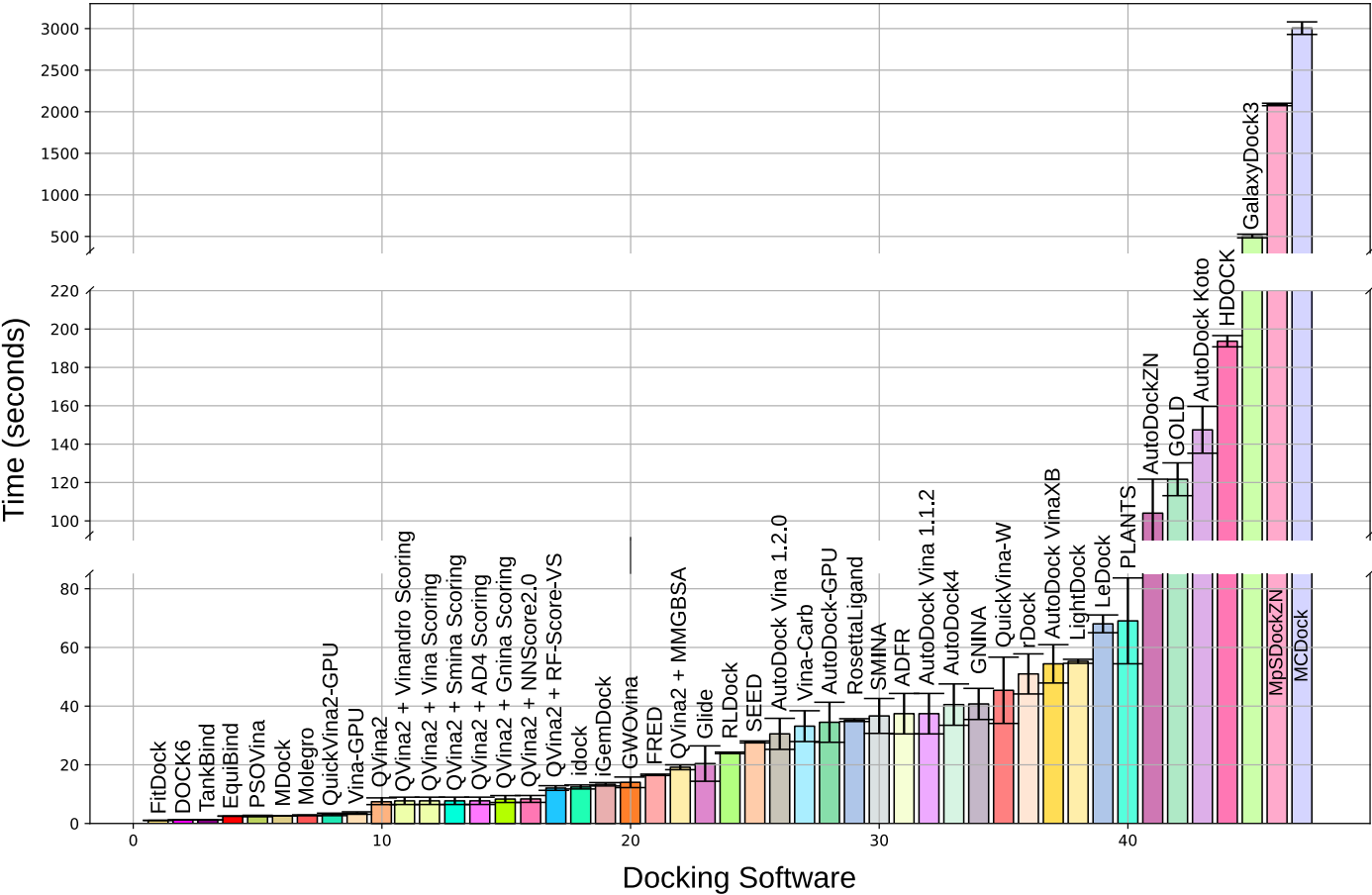
Extended Data Fig. 7 | Comparison of Adaptive Target-Guided (ATG) Primary Virtual Screens and standard ULVS. 100 million compounds were screened with standard ULVS of randomly selected compounds, as well as with ATGs where the 100 million compounds to be screened were selected based on the ATG prescreen. Violin plots of the docking scores of the top 50 virtual screening

hits obtained by standard ULVSs (blue) and ATG-VS (rose). Stripplots of the data points are overlaid. The docking scores of the top 50 compounds of the adaptive screens are considerably better for all target proteins than for the standard ULVSs.

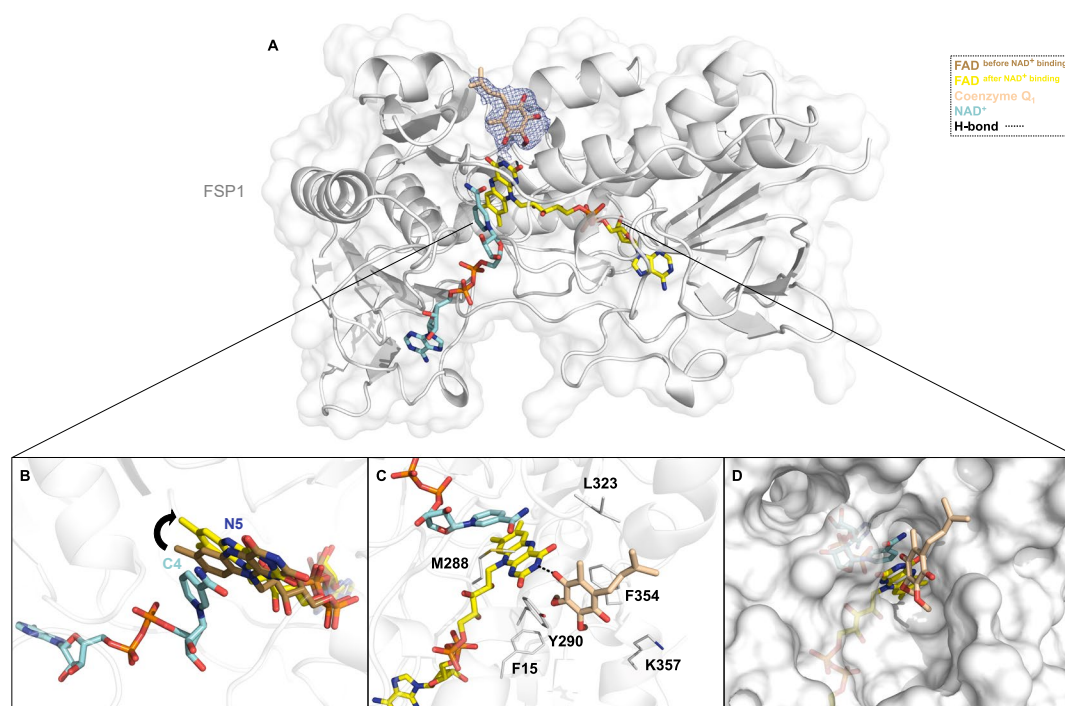


Extended Data Fig. 8 | Conceptual Overview of AdaptiveFlow Unity. AFU takes as input a user-specified string representation of a molecule and processes the provided molecule by desalting, neutralizing, and generating stereoisomer (with RDKit). Subsequently, based on the docking program choice, all ligands

are converted into 3D conformers (with OpenBabel) using a default protonation state of 7.4. Finally, built-in error checks are performed depending on a user's docking program choice, returning the docking pose and scoring function value. Steps in which the parameters of the user are involved are colored in white.



Extended Data Fig. 9 | Docking Time Estimates. Timing estimates for running a single calculation averaged over 25 independent runs for all docking programs. The calculations were run using AFU.



Extended Data Fig. 10 | Co-Crystal Structures of FSP1 in Complex with its Cofactors and Substrate. We determined three crystal structures of FSP1: (i) bound to FAD alone, (ii) bound to FAD and NAD⁺, and (iii) bound to FAD, NAD⁺, and coenzyme Q₁. **(a)** Overall structure of the fully assembled FSP1 complex with FAD (yellow carbons), NAD⁺ (cyan), and coenzyme Q₁ (beige). The polder omit map for coenzyme Q₁ is contoured at the 3.0 σ level. **(b)** Superposition of structures (i) and (ii) reveals that NAD⁺ binding induces a rotation of the FAD

isoalloxazine ring. The FAD conformation prior to NAD⁺ binding is shown in brown, and the conformation after NAD⁺ binding is shown in yellow (see also Supplementary Fig. 19a). **(c)** In the ternary complex, coenzyme Q₁ interacts with Tyr290 and Phe354 via edge-to-edge aromatic stacking and forms a hydrogen bond (dashed black line) with the flavin's pyrimidine ring. **(d)** NAD⁺ and coenzyme Q₁ occupy distinct but converging channels that meet at the FAD active site, illustrating the architecture of the substrate-binding cavity.

Corresponding author(s): Andrea Mattevi, Christoph Gorgulla, Haribabu Arthanari

Last updated by author(s): Ap 28, 2026

Reporting Summary

Nature Portfolio wishes to improve the reproducibility of the work that we publish. This form provides structure for consistency and transparency in reporting. For further information on Nature Portfolio policies, see our [Editorial Policies](#) and the [Editorial Policy Checklist](#).

Statistics

For all statistical analyses, confirm that the following items are present in the figure legend, table legend, main text, or Methods section.

n/a Confirmed

- ☐ ☒ The exact sample size (n) for each experimental group/condition, given as a discrete number and unit of measurement
- ☐ ☒ A statement on whether measurements were taken from distinct samples or whether the same sample was measured repeatedly
- ☒ ☐ The statistical test(s) used AND whether they are one- or two-sided
Only common tests should be described solely by name; describe more complex techniques in the Methods section.
- ☒ ☐ A description of all covariates tested
- ☒ ☐ A description of any assumptions or corrections, such as tests of normality and adjustment for multiple comparisons
- ☒ ☐ A full description of the statistical parameters including central tendency (e.g. means) or other basic estimates (e.g. regression coefficient) AND variation (e.g. standard deviation) or associated estimates of uncertainty (e.g. confidence intervals)
- ☒ ☐ For null hypothesis testing, the test statistic (e.g. F , t , r) with confidence intervals, effect sizes, degrees of freedom and P value noted
Give P values as exact values whenever suitable.
- ☒ ☐ For Bayesian analysis, information on the choice of priors and Markov chain Monte Carlo settings
- ☒ ☐ For hierarchical and complex designs, identification of the appropriate level for tests and full reporting of outcomes
- ☒ ☐ Estimates of effect sizes (e.g. Cohen's d , Pearson's r), indicating how they were calculated

Our web collection on [statistics for biologists](#) contains articles on many of the points above.

Software and code

Policy information about [availability of computer code](#)

Data collection

Autodock Vina (1.1.2 and 1.2), Smina, QuickVina 2, FitDock, DOCK6, TankBind, EquiBind, PSOVina, MDock, Molegro, FRED, Glide, RLDock, SEED, AutoDock-GPU, RosettaLigand, GNINA, Light Dock, LeDock, PLANTS, GOLD, HDock, COOT, Topspin, NMRPipe, and CCPN. The platform AdaptiveFlow developed for this manuscript is freely available on GitHub at <https://github.com/QuantumAI4Bio/JChemAxon's-JChem-Package>, Open Babel, and RDKit were used for ligand preparation. Biochemical and cellular-based assays were performed using a Cary 100 UV-Vis spectrophotometer and ClarioStar plate reader. X-ray diffraction data for protein crystallography were collected at the Swiss Light Source and European Synchrotron Radiation Facility in Grenoble (France). NMR of compounds: Bruker spectrometer equipped with a cryogenic probe and operating at a ^1H frequency of 400MHz. Data were acquired using TopSpin (Bruker, version 3.6.2)

Data analysis

Python; GraphPad Prism 9; For protein crystallography: data analysed and refined with CCP4 package, WinCoot CCP4 was used for electron density and model building

For manuscripts utilizing custom algorithms or software that are central to the research but not yet described in published literature, software must be made available to editors and reviewers. We strongly encourage code deposition in a community repository (e.g. GitHub). See the Nature Portfolio [guidelines for submitting code & software](#) for further information.

Data

Policy information about [availability of data](#)

All manuscripts must include a [data availability statement](#). This statement should provide the following information, where applicable:

- Accession codes, unique identifiers, or web links for publicly available datasets
- A description of any restrictions on data availability
- For clinical datasets or third party data, please ensure that the statement adheres to our [policy](#)

The FSP1 ancestral sequence generated in this study has been submitted for deposition in the GenBank database. The accession code is PQ619396. FSP1 structures have been deposited with PDB under accession codes: 9IFZ 9IFT; 9IFY; 9IFU, 9IG9, 9QRT. The structure of PARP1-inhibitor complex has been deposited with PDB code 8VYH. The ready-to-dock Enamine REAL Space can be explored online via an interactive web interface available on the AdaptiveFlow homepage <http://adaptive-flow.ai/>.

Code Availability

The source code of AdaptiveFlow is released under the GNU GPL v2.0, a free and open-source license. The code can be accessed via GitHub. The Adaptive Flow GitHub repos can be found at <https://github.com/LigandUniverse>. The AFLP and AFVS modules of AdaptiveFlow can be found in AFLP (<https://github.com/LigandUniverse/AFLP>) and the AFVS (<https://github.com/LigandUniverse/AFVS>) GitHub repositories, respectively. The modules AdaptiveFlow Unity and Adaptive Flow Unity Parallelized can be found in the GitHub repo named AFU (<https://github.com/LigandUniverse/AFU>) and AFUpart (<https://github.com/LigandUniverse/AFUpart>), respectively

Research involving human participants, their data, or biological material

Policy information about studies with [human participants or human data](#). See also policy information about [sex, gender \(identity/presentation\), and sexual orientation](#) and [race, ethnicity and racism](#).

Reporting on sex and gender

Reporting on race, ethnicity, or other socially relevant groupings

Population characteristics

Recruitment

Ethics oversight

Note that full information on the approval of the study protocol must also be provided in the manuscript.

Field-specific reporting

Please select the one below that is the best fit for your research. If you are not sure, read the appropriate sections before making your selection.

☒ Life sciences ☐ Behavioural & social sciences ☐ Ecological, evolutionary & environmental sciences

For a reference copy of the document with all sections, see [nature.com/documents/nr-reporting-summary-flat.pdf](https://www.nature.com/documents/nr-reporting-summary-flat.pdf)

Life sciences study design

All studies must disclose on these points even when the disclosure is negative.

Sample size

Data exclusions

Replication

Randomization

Blinding

Reporting for specific materials, systems and methods

We require information from authors about some types of materials, experimental systems and methods used in many studies. Here, indicate whether each material, system or method listed is relevant to your study. If you are not sure if a list item applies to your research, read the appropriate section before selecting a response.

Materials & experimental systems

n/a	Involved in the study
☐	☒ Antibodies
☐	☒ Eukaryotic cell lines
☒	☐ Palaeontology and archaeology
☒	☐ Animals and other organisms
☒	☐ Clinical data
☒	☐ Dual use research of concern
☒	☐ Plants

Methods

n/a	Involved in the study
☒	☐ ChIP-seq
☒	☐ Flow cytometry
☒	☐ MRI-based neuroimaging

Antibodies

Antibodies used	Antibodies used: anti-AIFM2 antibody (clone 1115 1129 ZooMAb, rabbit monoclonal, Sigma-Aldrich): HRP-conjugated anti-rabbit secondary antibody(Millipore); Protein Strep-Tactin HRP-conjugate(BioRad). anti FSP1 antibody (HY-P80679 from MCE), a Rabbit-derived; HRP conjugated anti Rabbit secondly antibody (Jackson ImmnoResearch Laboratories, 711-035-152); anti GPX4 antibody (67763-1-1g from Proteintech), a mouse-derived; HRP conjugated anti Mouse secondly antibody (Jackson ImmnoResearch Laboratories 715-035-150); anti GapDH antibody (6C5) (AM4300 from Invitrogen), a mouse- derived; HRP conjugated anti Mouse secondly antibody (Jackson ImmnoResearch Laboratories 715-035-150); anti-BRCA1 (MilliporeSigma, 07-434, Rabbit polyclonal), anti-Alpha/beta-Tubulin (CST, 2148S, Rabbit polyclonal), HRP-linked anti-rabbit secondary antibody (Cytiva, NA9341ML)
Validation	anti-AIFM2 antibody dilution 1:1000 HRP-conjugated anti-rabbit secondary antibody dilution 1:10000; Protein Strep-Tactin HRP-conjugate antibody dilution 1:10000; Anti FSP1 antibody dilution 1:1000; HRP conjugated anti Rabbit secondary antibody 1:10000; Anti GPX4 antibody dilution 1:500; HRP conjugated anti Mouse secondary antibody 1:10000; Anti GapDH antibody dilution 1:2000; HRP conjugated anti Mouse secondary antibody 1:10000; anti-BRCA1 dilution 1:1000, anti-Alpha/beta-Tubulin 1:1000, HRP-linked anti-rabbit secondary antibody 1:5000

Eukaryotic cell lines

Policy information about [cell lines and Sex and Gender in Research](#)

Cell line source(s)	MDA-MB-436 were purchased from ATCC (HTB-130); Jurkat cells were purchased from ATCC; HEK293-EBNA1-6f cells were obtained from Yves Durocher (NRC-BRI, Montreal, Canada; patent number US8551774).
Authentication	For MDA-MB-436: The cell line was authenticated by STR profiling. For Jurkat and HEK293-EBNA1-6f: These cell lines were authenticated using morphology and functional assays
Mycoplasma contamination	For MDA-MB-436: The cell line was regularly tested for mycoplasma contamination using a PCR-based assay, For Jurkat and HEK293-EBNA1-6f: : The were not tested for Mycoplasma contamination.
Commonly misidentified lines (See ICLAC register)	N/A

Plants

Seed stocks	N/A
Novel plant genotypes	N/A
Authentication	N/A