

Adaptive Sampling Methods for Molecular Dynamics in the Era of Machine Learning

Diego E. Kleiman, Hassan Nadeem, and Diwakar Shukla*



Cite This: *J. Phys. Chem. B* 2023, 127, 10669–10681



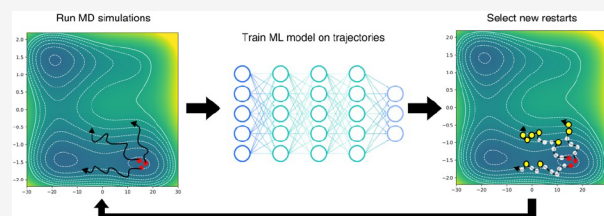
Read Online

ACCESS |

Metrics & More

Article Recommendations

ABSTRACT: Molecular dynamics (MD) simulations are fundamental computational tools for the study of proteins and their free energy landscapes. However, sampling protein conformational changes through MD simulations is challenging due to the relatively long time scales of these processes. Many enhanced sampling approaches have emerged to tackle this problem, including biased sampling and path-sampling methods. In this Perspective, we focus on adaptive sampling algorithms. These techniques differ from other approaches because the thermodynamic ensemble is preserved and the sampling is enhanced solely by restarting MD trajectories at particularly chosen seeds rather than introducing biasing forces. We begin our treatment with an overview of theoretically transparent methods, where we discuss principles and guidelines for adaptive sampling. Then, we present a brief summary of select methods that have been applied to realistic systems in the past. Finally, we discuss recent advances in adaptive sampling methodology powered by deep learning techniques, as well as their shortcomings.



INTRODUCTION

Computer simulation studies have been invaluable tools to study atomic scale phenomena. Experimental observations and theoretical predictions can, in theory, be validated through these simulations. Molecular dynamics (MD) simulations are a powerful technique that can probe these molecular systems at atomic scales. MD simulations iteratively solve equations of motion, which allows the molecular system to evolve over time steps at the order of 1–2 fs and perform sampling to recover statistical ensembles. Although MD simulations offer unparalleled insight into the atomic world, there are many limitations to this approach.

MD simulations require the interaction potential to be defined in terms of a force field, which in turn makes the simulation accuracy dependent upon this choice of parameters. Hence the simulations will not offer desired insights in a general sense, rather only for the application for which the force field has been parametrized. In this regard many force fields have been proposed like CHARMM,¹ AMBER,² GROMOS,³ OPLS,⁴ etc., for simulation of biological as well as materials systems.

Another major limitation that MD simulations suffer from is the time scale problem. MD simulations generally employ an integration time step of 1–2 fs, corresponding to the fastest degree of freedom for the molecular system under study. Many processes of practical interest, especially biological processes like protein folding, ligand binding, etc., are of the order of milliseconds or even higher. For these processes, the sampling probability decays exponentially with energy (Boltzmann distribution), so high energy or rare transitions pose a

challenge. In other words, due to the decaying probability of sampling a high energy state, the time that it takes a simulation to cross an energy barrier increases exponentially with the size of the barrier, and this gives origin to the time scale problem in MD. A traditional long MD simulation can remain stuck in a metastable basin and fail to sample the conformational landscape as desired. A nonspecialized computer provides a computational speed of the order of nanoseconds per day which would require years of computation to reach the millisecond stage. To deal with this bottleneck, many alternate approaches have been investigated.

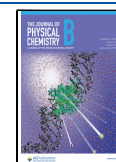
Coarse-graining of the system under study has been a popular approach to studying such biological processes. In coarse-graining, sets of atoms are collectively represented by “beads” which act as a representative “pseudo-atom” that hopes to capture the chemical behavior of the modeled group of atoms; e.g., MARTINI⁵ is a commonly used coarse-grained force field to this end. This clustering reduces the number of atoms (reducing the number of motion equations to be solved) hence reducing the computational expense as complexity for MD simulations is $O(N \log N)$,⁶ for N number of atoms. Another speedup comes from the idea that, because finer

Received: July 18, 2023

Revised: November 14, 2023

Accepted: November 16, 2023

Published: December 11, 2023



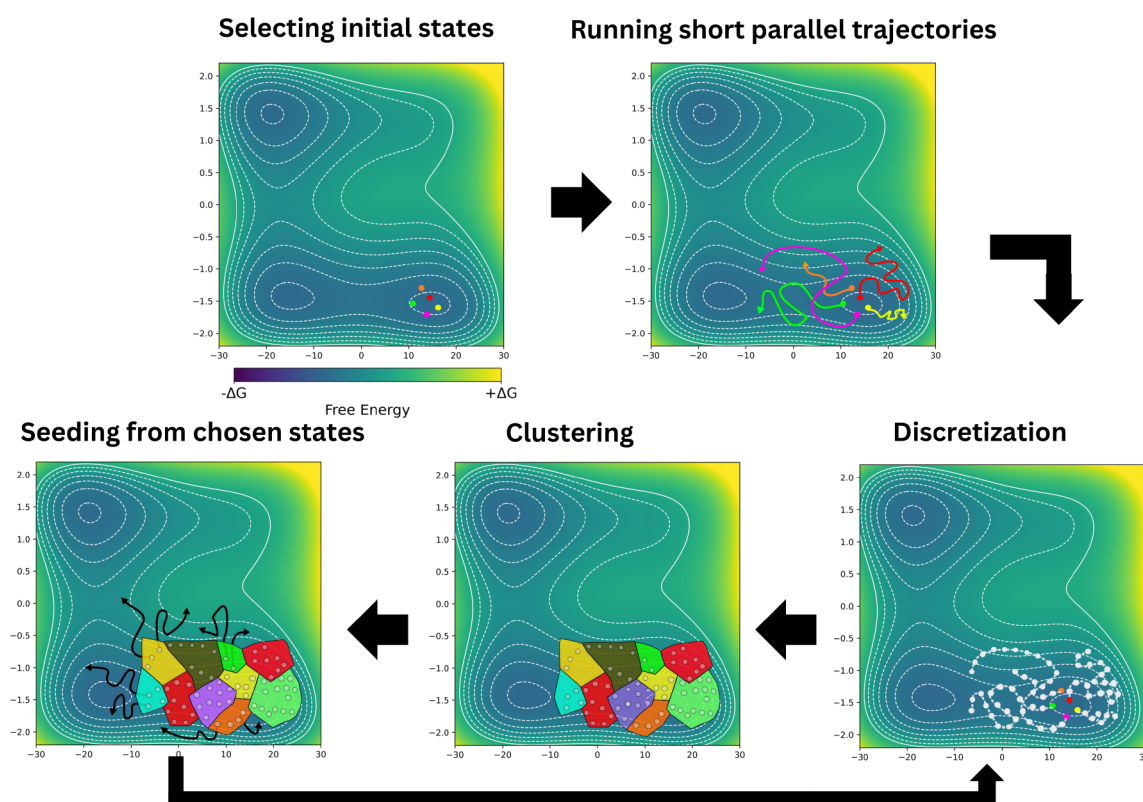


Figure 1. Adaptive sampling, starting from initial states and running short simulations, discretization of conformations, clustering into representative states, and reseeding from states chosen via the chosen scheme.

degrees of freedom have been coarse-grained, the integration time step could now be increased from 1–2 fs to 20–40 fs. Although both of these speedups are significant, there are some limitations⁷ to this approach as well. The outcome of a coarse-grained simulation is heavily dependent upon the choice of coarse-graining scheme; hence, there is a possibility of coarse-graining out potentially important degrees of freedom for our system. Another caveat is that the energy surface gets smoothed out due to this coarse-graining effect; therefore, the simulation time does not equal the actual time. In response, calculations of the dynamic quantities must be scaled accordingly.

Enhanced sampling methods such as adaptive sampling methods, replica exchange methods, localization methods, biasing methods (adaptive and nonadaptive), and more have been proposed as a strong substitute to address this sampling dilemma. The number and types of enhanced sampling methods that have been proposed are too many to mention here, and the reader is directed to this excellent review.⁸ Although these methods perform well for specific systems, there are still some drawbacks that limit their applicability in a general sense. For example, biased enhanced sampling methods (e.g., metadynamics) add an external bias to the system. This perturbation modifies the underlying potential energy surface causing the system to lose kinetic information, which can only be recovered in theory by nontrivial methods.^{9,10} Additionally, these methods could potentially sample unphysical conformations due to external forces. Other techniques (e.g., replica-exchange¹¹) work well for enthalpic barriers but perform relatively poorly for systems where entropic barriers are dominant.

The class of methods addressed in this Article can generally be described as unbiased adaptive sampling methods. The key idea is to, after an initial run of short trajectories, strategically restart these trajectories based on some criterion. In adaptive sampling methods, instead of conventional long MD simulations, multiple short simulations are run in parallel. Then, states from the resultant trajectories are selected according to some criterion to run the next round of simulations. It is the choice of this adaptive selection criterion that distinguishes the different types of methods in this class. The process is illustrated in Figure 1.

The intuition behind this adaptive sampling is to start sampling from states that have been relatively less sampled and which would be more likely to overcome free-energy barriers in rare events like protein folding. These short trajectories are then “stitched” together, often by using Markov state models, where states are clustered and probabilities of state transformations are recorded in a transition matrix and analyzed. Time-lagged independent component analysis (tICA) is a linear transformation method that can be used for dimensionality reduction and is often¹² applied to the time-lagged correlation matrices of the features in the MSM construction pipeline.

In this Perspective, we outline the theoretical underpinnings of adaptive sampling as well as major methods that have been developed recently. An attempt has been made to categorize these techniques into three classes. The first categorization is termed Markov state model (MSM) inspired adaptive sampling and outlines the sampling methods that make use of the Markovian assumption to recover thermodynamic quantities of interest as well as kinetic information from multiple short trajectories. The next class is named reinforce-

ment-learning-inspired adaptive sampling methods. As the name suggests, these methods draw inspiration from reinforcement learning, where exploration and/or exploitation in the sampling space are favored/penalized via an objective function. Beyond these, we also discuss other methodologies such as weighted ensemble techniques where the configuration space is divided into bins and walkers (trajectories) in the bins are replicated/terminated according to their favorability toward the target state.

THEORY

The practice and theory of adaptive sampling methods have developed rather unevenly over their history. For this reason, the theoretical characterization of certain methods is more advanced than that for others. Nonetheless, due to the complicated statistical behavior of high-dimensional dynamical systems, simplifying assumptions of varying strength are applied in the derivation of theoretical principles (e.g., expected advantage over naive methods) and principled guidelines (e.g., optimal allocation strategies). The validity of these assumptions is not usually guaranteed in MD simulations. But before exploring these theoretical principles and their underlying assumptions, it is useful to first define the quantities that adaptive sampling methods intend to estimate, as these will motivate the theory.

Adaptive sampling methods are focused on accelerating the sampling of state transitions and the convergence of thermodynamic and kinetic models of the molecular system under study.^{8,13–15} The free energy landscape encodes the thermodynamics of a molecular system because it provides the probability of observing a conformation under the simulated thermodynamic ensemble, $P(\mathbf{x}) = \frac{1}{Z} e^{\beta F(\mathbf{x})}$, where $\mathbf{x}$ is a molecular conformation, F is the free energy, Z is the canonical partition function, and β is the thermodynamic beta.

A kinetic model is a mathematical model that describes the time evolution of a system. In the context of molecular dynamics, a kinetic model is composed of a set of state definitions (generally expressed as boundaries in conformational space) and the average rates of change or mean first passage times (MFPTs) between the states. The MFPT can be defined as the average number of trajectory steps that it takes to reach one state from another. The Hill relation can be used to compute the MFPT from trajectory data¹⁶

$$\langle T_B \rangle_A = \left(\frac{d}{dt} \langle N_t \rangle_\pi \right)^{-1} \quad (1)$$

where $\langle T_B \rangle_A$ is the MFPT from source state A to target state B and $\langle N_t \rangle_\pi$ is the number of arrivals on state B at time t given the steady state distribution π . This estimate of the MFPT is correct only under the assumptions of steady state convergence, which is computationally challenging for complex systems, and recycling boundary conditions (trajectories that reach B are immediately restarted from A). In other words, state B is a “sink” that absorbs all arriving trajectories and state A is a “source” where all trajectories originate. The distribution π is achieved once the rate of arrival to state B (the sink) does not change as time increases. A mathematically rigorous treatment of the Hill relation under molecular dynamics is available in the literature.¹⁷ MFPTs are central to the characterization of molecular systems because, once known, they allow us to calculate other observables through kinetic modeling.¹⁸ Depending on the adaptive sampling method

utilized, the measured MFPTs might be biased. For example, adaptive sampling results in statistically biased MFPTs, while the weighted ensemble takes care of such bias on the fly by assigning weights to trajectories. Methods that produce biased MFPTs employ posthoc statistical models, like Markov state models (MSMs)¹⁵ or generalized master-equation-based models (GMEMs),^{19,20} to recover the unbiased MFPTs.

Adaptive sampling can help accelerate the sampling of transition states and kinetic models via two mechanisms, neither of which is exclusive to adaptive sampling methods: trajectory parallelization and selective seeding. Trajectory parallelization refers to running many unbiased trajectories simultaneously. On the other hand, selective seeding is the act of restarting a simulation from a set of specific configurations chosen according to a criterion.

To analyze the theoretical contribution of each mechanism or, at least, under what circumstances each mechanism is helpful, we can first remove the “seeding contribution” from the equation and analyze the advantage of using parallel trajectories only. Interestingly, it is an unrelated field (stochastic resetting) that sheds light on this question. In stochastic resetting the main premise is that a diffusion process (the molecular dynamics simulation) is set back to its original position (conformation) after some random number of time steps.²¹ Clearly, there is no selective seeding under stochastic resetting because the system is always set back to the same seed. We note that, in stochastic resetting, there is no actual parallelization of trajectories either. However, this is an implementation detail. If they were run in parallel, one could sample permutations from the set (akin to bootstrapping) to recover “ordered” trajectories and the statistical analysis holds.

Under stochastic resetting we must consider two random variables: T and R .²² T is the number of time steps that it takes for an individual trajectory started in state A to reach the target state B without being restarted. R is the maximum possible length of the individual trajectory (after R time steps, it is restarted). If $T < R$, the FPT under resetting (T_r) is measured as T . However, if $R \leq T$, then R is added to T_r and a new trajectory is sampled. Therefore, we can express T_r in a recursive fashion, $T_r = \min(T, R) + \mathbf{1}_{R \leq T} T'_r$, where T'_r is i.i.d. to T_r and $\mathbf{1}_{R \leq T}$ is an indicator function that halts the summation once $T < R$.²² Note that, in the case of the parallel simulations, we would need to “discard” any trajectories in the permutation that come after one where $T < R$. Taking the expectation of this expression, we get²²

$$\langle T_r \rangle = \frac{\langle \min(T, R) \rangle}{P(T < R)} \quad (2)$$

This is the “effective” MFPT measured by the restarted (or short and parallel) trajectories.

Although the output of the expression depends on the specific probability distributions of T and R , it tends to be smaller than the unbiased MFPT ($\langle T_B \rangle_A$) because the distribution of the FPT can be heavy-tailed, a feature that is countered by the restarting procedure. In fact, as it has been proven^{22,23} and has been empirically explored,²⁴ if the coefficient of variation (standard deviation divided by the mean) of the unbiased FPT is greater than 1, then $\langle T_r \rangle < \langle T_B \rangle_A$. This translates into a speedup in the sampling of state transitions. For example, researchers have shown a toy system where the speedup reaches 1 order of magnitude with a coefficient of variation of 2.9.²⁴ It must be noted that the actual

speedup depends on the distribution of the unbiased FPT, not only on the coefficient of variation.²⁵ A method to recover the unbiased MFPT from the restarted one was proposed,²⁴ but it seems that the theoretical bound for the error remains elusive. Nonetheless, other statistical methods, such as MSMs¹⁵ and GMEMs,^{19,20} might prove useful to recover the unbiased MFPT from restarted trajectories.

The takeaway message from this analysis is that, even without any selective seeding, one could have a considerable speedup in state transition sampling and MFPT convergence from parallelization only. This would translate into fewer simulation time steps required to reach the target state and a wall-clock time saved to compute the MFPT. Therefore, when testing new adaptive sampling methods, it is important to include a sensible baseline that accounts for the parallelization advantage. For example, one could compare the proposed technique to another parallel method. If one merely claims that an adaptive sampling method improves upon long, continuous MD simulations without explicitly providing the distribution of the unbiased FPT or its coefficient of variation, then it remains ambiguous whether the seeding strategy is actually responsible for the speedup.

Now that we have considered the situation where parallelization alone provides a speedup, we turn to theoretical results pertaining to the selective seeding advantage. We will focus on three specific approaches based on the theoretical clarity that they bring into the discussion, noting that they are not routinely used in practice and their underlying assumptions are demanding in most cases.

The first one is termed coupled parallel trajectories.²⁶ In this method, M parallel trajectories are started from the same state. After a single trajectory has sampled a state transition, all of the trajectories are stopped, and M new trajectories are started in the new state. It is assumed that a transition can occur between any two states (all states are connected). By making the additional simplifying assumption that the FPTs between states are exponentially distributed, then we can draw explicit results. For this, note that the problem setup can be expressed as a system of differential equations, $\frac{d}{dt}\mathbf{p} = K\mathbf{p}$, where K is a matrix that contains the transition rates and $\mathbf{p}$ contains the state populations. The long time solution to this system is $\mathbf{p} = \sum_i c_i \mathbf{v}_i e^{\lambda_i t}$, where λ_i and $\mathbf{v}_i$ are the i th eigenvalue and eigenvector of K and c_i 's are constants that depend on boundary conditions. Further assuming Markovianity and absorbing boundary conditions (the simulations are stopped at the target state), we get that the MFPT to the target state is

$$\langle T'_B \rangle_A = \sum_i \frac{c_i v_{iB}}{M \lambda_i} = \frac{1}{M} \langle T_B \rangle_A \quad (3)$$

where $\langle T'_B \rangle_A$ is the MFPT estimated from the coupled trajectories and v_{iB} is the B th component of the i th eigenvector. A full derivation of the second equality is available in the original study.²⁶ The key result is that, in some simple cases (i.e., systems with fully connected states and exponential transition rates), a greedy strategy where all trajectories are moved to the newly discovered states as soon as they are seen can provide a linear sampling speedup. Of course, realistic molecular systems are generally too complex for the assumptions to hold; they include unconnected states, kinetic traps, and other features that frustrate this type of greedy scheme. Moreover, when comparing against M parallel but uncoupled trajectories the relative speedup is not linear.²⁶

The next approach that we will consider offers a more sophisticated view, since it considers selective seeding as a means to reduce the variance in kinetic models rather than simply increasing the sampling rates.²⁷ The idea behind this method is to perform error analysis on the first nontrivial eigenvalue of the Markovian transition probability matrix and then selectively seed new simulations from the state with the largest contribution to its variance. The result is a sampling scheme that improves the resolution of the slowest relaxation process. To achieve this, two key assumptions are made: (1) the transition probabilities converge to normal distributions after enough transitions have been sampled and (2) the first-order Taylor expansion around the eigenvalue is a good enough approximation of the effect of small perturbations in the transition probability matrix. To express the selection criteria mathematically, it is useful to define certain quantities first. Let $\bar{\mathbf{k}}_i = K_i$ be the normalized vector containing all transition probabilities for starting state i , and let $\mathbf{s}_i^j = \nabla_{\mathbf{k}_i} \lambda |_{\bar{\mathbf{k}}_i}$. Here, $\nabla_{\mathbf{k}_i}$ is the gradient operator over the variables in $\mathbf{k}_i$ (all transition probabilities originating from state i), so $\nabla_{\mathbf{k}_i} \lambda |_{\bar{\mathbf{k}}_i}$ is the gradient of the eigenvalue evaluated at $\mathbf{k}_i = \bar{\mathbf{k}}_i$. This gradient is termed the sensitivity vector, which linearly approximates how much the eigenvalue λ will vary with changing an element in the probability transition matrix. These values are used to compute $\bar{q}_i = (\mathbf{s}_i^j)^T [\text{diag}(\bar{\mathbf{k}}_i) - \bar{\mathbf{k}}_i \bar{\mathbf{k}}_i^T] \mathbf{s}_i^j$. Then, the variance of the eigenvalue is given by $\sigma^2 = \sum_i \bar{q}_i / (w_i + 1)$, where w_i is a normalization factor for the transition counts from state i .²⁷ If we add m new samples to state i (assuming the transition probabilities stay constant), then the state that will result in the highest reduction in the variance will be given by $i = \text{argmax}(\bar{q}_i / (w_i + 1) - \bar{q}_i / (w_i + m + 1))$.²⁷ When this selection criterion is used, the variance of the first nontrivial eigenvalue decays faster than when using other forms of parallel sampling.²⁷ This work shows that, given a series of approximations, one can reach an elegant, closed-form solution that determines the criterion for adaptive sampling. Follow-up works showed that, in practice, other adaptive sampling schemes perform better in terms of error reduction,²⁸ a sign that the assumptions used in this analysis (e.g., linearity of eigenvalue perturbations) are too stringent.

The last approach that will be described differs from the previous ones because it was formulated to work under weighted ensembles rather than Markov state models.¹⁶ Similar to the approach described before, the goal is to reduce the variance in a metric of importance to the kinetic model. In this case, rather than minimizing the variance of the first nontrivial eigenvalue, the goal is to reduce the variance in the estimated MFPT from the source state to a sink state under recycling boundary conditions (trajectories that enter the sink are immediately restarted from the source). Since weighted ensemble simulations involve stopping unproductive trajectories ("merging") and allocating productive ones ("splitting"), the idea behind this approach is to perform these actions following optimal coordinates that guarantee that the variance of the MFPT will be reduced. For this reason, two optimal coordinates must be defined, the flux discrepancy function, $h(\mathbf{x})$, and the flux variance function, $\nu(\mathbf{x})^2$,

$$h(\mathbf{x}) = \frac{\langle T_B \rangle_{\pi} - \langle T_B \rangle_{\mathbf{x}}}{\langle T_B \rangle_A} \quad (4)$$

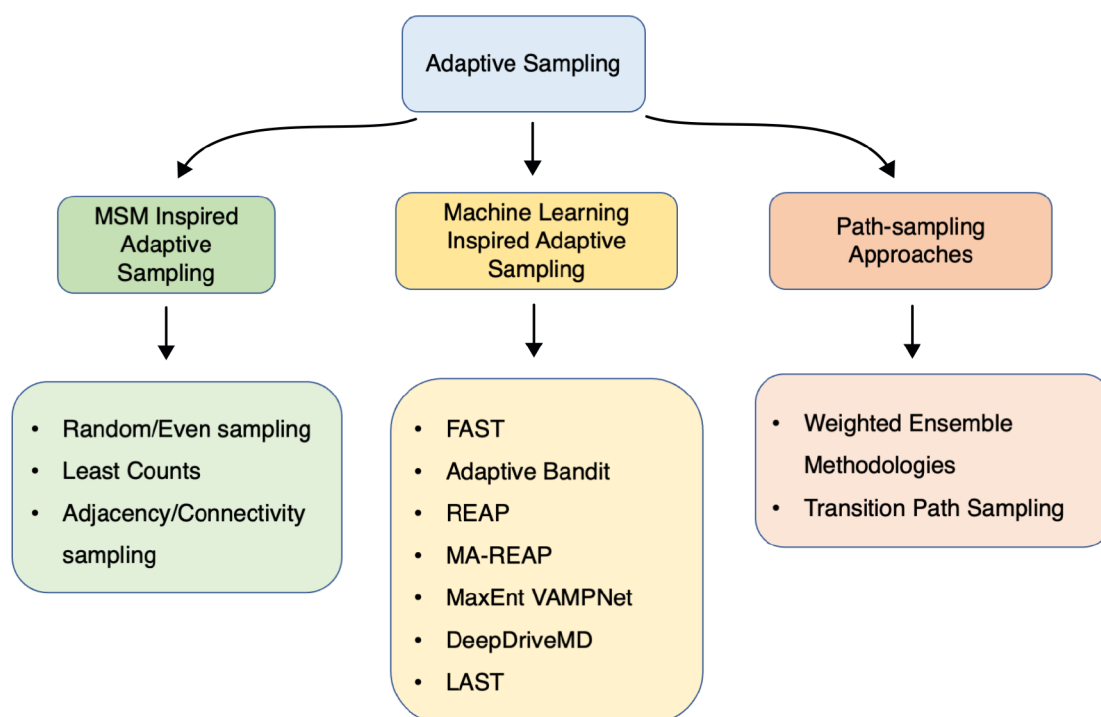


Figure 2. Hierarchy of adaptive sampling for enhanced sampling in MD simulations.

$$v(\mathbf{x})^2 = \frac{1}{\tau} \text{Var}_{\mathbf{x}}[\mathbf{1}_{X_T \in B} + h(X_T)] \quad (5)$$

where $\mathbf{x}$ is a conformation, $\langle T_B \rangle_{\pi}$ is the MFPT to the target state from the steady state distribution, $\langle T_B \rangle_{\mathbf{x}}$ is the MFPT from location $\mathbf{x}$, and $\langle T_B \rangle_A$ is the MFPT from the source state A (which, unfortunately, is the value we were trying to estimate in the first place). Furthermore, τ is a time interval over which the variance is computed, X_T refers to a trajectory given by the Markovian dynamics of the system, and $\mathbf{1}_{X_T \in B}$ is an indicator function that determines whether $\mathbf{x}$ is in state B at each step in the trajectory. In other words, $h(\mathbf{x})$ can be interpreted as the normalized kinetic distance between a point in phase space and the steady state distribution, while $v(\mathbf{x})^2$ gives us the expected change in flux into B given a trajectory started at $\mathbf{x}$. Of course, it is quite contradictory that we need to know the value that we want to estimate, $\langle T_B \rangle_A$, to compute the optimal coordinates. We also require the steady state distribution $\pi(\mathbf{x})$, which is typically challenging to compute for complex systems. Nonetheless, the authors of the original study propose to estimate the necessary values between simulation rounds by fitting MSMs with the available data and using the estimates to approximate the optimal coordinates.¹⁶ These coordinates are then used with a binning strategy to define regions of phase space which must be allocated the same number of trajectories to optimally reduce the variance in the measured MFPT. Namely, given H bins, we must set end points $h_0 < h_1 < \dots < h_H$ such that¹⁶

$$\int_{h_i \leq h(\mathbf{x}) \leq h_{i+1}} \pi(\mathbf{x}) v(\mathbf{x}) \, d\mathbf{x} = \text{constant} \quad (6)$$

In simpler terms, this means that we should weight a position in phase space not only by the density of the steady state distribution, $\pi(\mathbf{x})$, but also by its contribution to the MFPT fluctuation, $v(\mathbf{x})$.

Following this allocation rule results in a minimization of the variance of the flux (and therefore the variance on the estimated MFPT). In comparison with the same number of “brute-force” parallel trajectories, we obtain a maximum ratio of variance reduction given by¹⁶

$$\frac{\text{Var}(J_{\text{BF}})}{\text{Var}(J_{\text{WE}})} = \frac{\int v(\mathbf{x})^2 \pi(\mathbf{x}) \, d\mathbf{x}}{\left(\int v(\mathbf{x}) \pi(\mathbf{x}) \, d\mathbf{x} \right)^2} \quad (7)$$

where $\text{Var}(J_{\text{BF}})$ is the variance in the flux measured from parallel trajectories and $\text{Var}(J_{\text{WE}})$ is the variance in the flux measured by the optimally allocated trajectories. Although this expression is not transparent for high-dimensional systems, it can be used to compute the optimal advantage attainable for weighted ensemble simulations of analytical potentials. In particular, this result was used to show that in the low temperature limit the advantage with respect to brute-force simulations is exponential in the largest energy barrier.¹⁶

To summarize, in this section, we introduced four theoretical approaches to dissect the advantage of adaptive sampling methods. As we moved from simple parallelization to sophisticated selective seeding criteria, we discussed the known theoretical results associated with each approach and their underlying assumptions. Notwithstanding the mathematical transparency afforded by these approaches, different methods have shown better performance in practice. In the next section, we discuss a wide variety of adaptive sampling methods that have also been applied to realistic systems. After this, we discuss recent advances in the field with a particular focus on machine-learning-based methods.

METHODS

Adaptive sampling (AS) methods can be generally categorized into Markov state model inspired methods and reinforcement learning inspired schemes. Both approaches have been

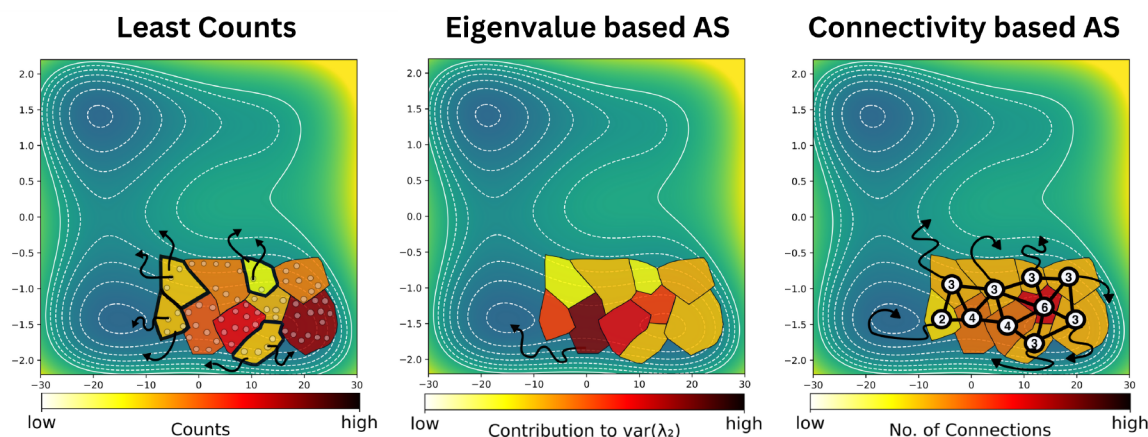


Figure 3. *Least Counts*: Least visited states; states with the lowest amount of counts are selected for reseeding. *Eigenvalue based AS*: States contributing most to the first nontrivial eigenvalue of T_{ij} are selected. *Connectivity based AS*: Least connected states; states having the lowest counts in the adjacency matrix are chosen.

extensively employed to investigate biophysical processes such as protein folding,²⁹ protein conformational changes,^{30–32} protein–ligand binding,^{33–35} membrane transport,^{36–38} and protein–protein association.³⁹ We describe the prominent methodologies in each of these categories followed by a short discussion on other methods, such as weighted ensemble,⁴⁰ transition path sampling,⁴¹ etc. Each of the methods is described briefly to give an overall view of the different adaptive sampling methodologies; readers are referred to relevant literature for an in-depth explanation of respective methods. Figure 2 illustrates this broad categorization.

MSM-Inspired Adaptive Sampling. The first class of methods that we describe are Markov-state-model-inspired adaptive sampling methods. The idea is to initialize simulations to explore and exploit regions of interest by using Markov state models as tools of analysis. This is done by seeding a set of short MD simulations; the resulting states are then clustered according to a kinetic/geometric criterion, and then out of these clusters, new states are chosen for seeding according to some methodology. It is this choice that distinguishes the method. The most commonly employed MSM-inspired AS methods are shown in Figure 3. In comparison with other adaptive sampling methodologies (e.g., metadynamics), where biased potentials (usually multidimensional Gaussians) are added to aid rare state transitions, MSMs have the advantage of recovering kinetics along with the thermodynamics by spectral analysis of the transition probability matrix.

One of the first methods was inspired by uncertainty analysis in MSMs.²⁷ Closed form expressions for the distribution of eigenvectors and eigenvalues of the transition matrix were derived. Correspondingly these distributions can be decomposed to calculate the contribution of variance to the first nontrivial eigenvalue of the Markovian transition probability matrix T_{ij} for each state and thereby seeding can be selected for states which contribute the maximum to this contribution. The method was shown to significantly increase precision for the villin headpiece, but the gain in precision was shown to be linked to the number of states the system had been coarse-grained into.

A conceptually simpler approach is to do a random⁴² selection of states to seed from. This method was shown to improve on generalized ensemble methods that fail to overcome entropic barriers at low temperatures and performed better for sampling of a hairpin folding system where

conformational changes are diffusion controlled. Variants of these two methods have also been employed. For instance, uniform (even) sampling can be done by seeding equally from states. Also, new simulations can be distributed among states in contribution to the uncertainty at the slowest rate.

Count based sampling²⁸ is another approach focused on sampling exploration. As the name suggests, states with a minimum number of counts are selected; in other words, states that have been less explored in the sampling are preferred for seeding to focus on exploration.

An adjacency based sampling²⁸ scheme has also been proposed. The idea is to start new simulations based on a connectivity based criterion. Because MSMs can be approximated generally as being similar to Cayley's trees topologically, the distant states in such a topological structure would be least visited; this is because MSMs can be thought of as network models for transition between states. So states with the least number of connections would be chosen.

Reinforcement-Learning-Inspired Adaptive Sampling. The next class of methods that we will discuss incorporates reinforcement learning ideas to sample the free-energy landscape. This can be done by following a gradient along a known property of interest and/or using a reward/penalty scheme that penalizes the system moving away from the target state and rewards the converse.

The choice of states to seed from can also be made on a ranking criterion where states are ordered according to a given metric significant for the application at hand. For instance this approach⁴³ has been applied for a trypsin–benzamide binding system and shown to improve upon traditional high-throughput experiments by an order of magnitude, where the ranking criteria used was mean residence time. A similar novel approach⁴⁴ was applied on sampling of pathways for rare conformational transitions, where the choice of states was based upon a distance metric between evolutionarily coupled residues.

Extending this idea of choosing via a ranking scheme, states can also be chosen according to a reward based scheme that favors those states which optimize some property of significance, e.g., RMSD in a protein folding system. Such a scheme named fluctuation amplification of specific traits (FAST)⁴⁵ was proposed and shown to be a significant improvement upon nondirected approaches mentioned above. The reward function that the algorithm seeks to

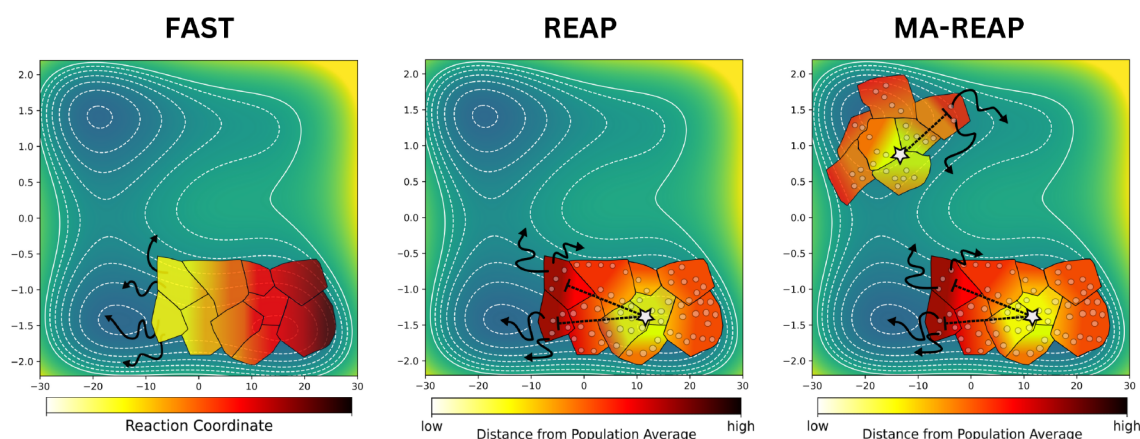


Figure 4. *FAST*: States that maximize movement along the gradient of interest are chosen for reseeding. *REAP*: States furthest away from the population average (and least visited) are chosen. *MA-REAP*: In multiagent REAP there are multiple agents that share information and drive the adaptive sampling process after each round.

optimize balances the exploration/exploitation trade-off. *FAST* does this by adjusting a control parameter that weighs the relative significance of the directed (exploitation of sampled space) and undirected (exploration of less sampled states) components. This was the first method to formally address the exploration/exploitation dilemma of sampling compared to purely exploratory schemes like least counts. The test case used was folding of villin protein among others. The main idea is that the system of interest will follow an approximately monotonic gradient of some property of interest. For example, in a protein folding problem, the transition from an unfolded state to a completely folded state can be thought of as a gradient of solvent accessible surface area which monotonically decreases as the protein folds itself, which is in essence optimizing along a single reaction coordinate (see Figure 4).

The multiarmed bandit problem is a famous problem in combinatorial optimization where an agent faces the choice of a policy that drives an action to return the maximum cumulative reward. AdaptiveBandit⁴⁶ is a sampling algorithm inspired by this problem and aims to maximize the reward, where the reward is set to be the mean of minus free energies of the conformation visited as a result of this action. The inspiration for this reward definition is inspired by the fact that in most MD simulations the aim is to find metastable states, for example, in protein folding problems, ligand binding, etc. The policy has to cater for the exploration–exploitation dilemma, which is at the heart of this optimization scheme, where the actions have to be exploratory to sample their unknown rewards and also exploitative to achieve maximum reward from the known best-rewarding space. In the MD perspective this translates to exploring the space of states which is less sampled (exploration) but also focusing on sampling the states which are known to drive the system toward the target end-state (exploitation).

Inspired by this exploration–exploitation dilemma of configuration states in MD simulations, a novel method called REAP⁴⁷ (reinforcement-learning-based adaptive sampling) was proposed. Conceptually, REAP is an extension of the counts method where preference for seeding is given to least visited states; however, in REAP this choice is based upon the reward function. This method differs and improves upon previously described directional methods such as *FAST* and Adaptive-Bandit (AB). In *FAST*, the choice of states follows a gradient along a property of interest (collective variable), and in AB, the

reward scheme minimizes the free energy of the configurations. In a more generic problem-agnostic framework, these approaches may not always give the best results. For some systems the particular collective variables (CVs, also called reaction coordinates) or the property of interest may not be known, and instead there may be a set of possibly relevant CVs available. In this scenario, a unidirectional gradient approach or a single objective reward strategy will not work. The problem then is to identify the relative importance of CVs as well as to understand that this importance (translates to weights in the algorithm) may be changing as the system progresses in the potential energy landscape. REAP solves this problem by dynamically computing these weights so that the system is driven along the CVs which contribute the most to the system moving toward the target state, essentially choosing states which are most distant from the average of all conformations in the adaptive sampling round (see Figure 4). A limitation to this algorithm that is addressed in MA-REAP (explained next) is that, if sampling is started from different states, then pooling the information can result in the rewarding scheme favoring only the states which drive the CVs to have extreme values.⁴⁸

This idea of reinforcement-learning-based adaptive sampling has further been extended to an algorithm called MA-REAP⁴⁸ (MA for multiagent). The core of the algorithm resembles REAP closely, and the addition has been the idea of having multiple agents that drive the adaptive sampling process after each round. Compared to the single agent traditional REAP algorithm, MA-REAP proposes that there may be multiple agents that share data. In REAP, the action space or the sampled configurations go through a clustering process; in MA-REAP, this is extended in the sense that all agents have a stake in each cluster for an action (choice of reseeding) according to the number of configurations present in that cluster because of that action from that agent. In this way, the reward function has a distributed essence according to each agent, allowing each agent to sample along an independent CV and sharing information only when in proximity of other agent states. This variation has shown to be an improvement over the traditional REAP and was shown to outperform previously introduced directional methods like *FAST* and count-based methods as well.

Other Approaches. In this section, we discuss other powerful approaches for enhanced sampling of biomolecular systems, which do not fall under the previous two classes.

Weighted ensemble^{14,40} (WE) methods are another class of methods designed to explore pathways for rare states. The basic idea behind weighted ensemble methods in MD simulations is to divide the system into multiple copies, each with a different set of initial conditions. These copies then evolve independently in parallel, and the idea is to replicate simulations that are favoring progression toward the target state while terminating the others. This concept can be thought to be based on this early idea of a splitting strategy: “When the sampled particle goes from a less important to a more important region, it is split into two independent particles, each one-half the weight of the original.”⁴⁹ So in essence weighted ensemble techniques involve splitting and merging trajectories based on their importance for a sampling criterion. The trajectories are assigned weights, and splitting decreases the weight, while merging increases it where the objective is to obtain unbiased observables statistically. Often, sampling in WE is performed using bins that are divisions of the configuration space, and the algorithm ensures the number of walkers (trajectories) per bin to be the same. In WE, transition rates can be calculated using the Hill relation without relying on a Markovian approximation by measuring the flux into the target state from the weights of the incoming trajectories. It is also possible to build MSMs with WE data taking the trajectory weights into account, as opposed to assigning equal weights to all observed transitions.⁵⁰

Nonetheless, WE presents limitations, as well. As the number of bins has an exponential dependence on the number of CVs, this method can incur significant computational cost. Variants of WE have been proposed that address this and other issues with a traditional approach. WExplore⁵¹ is an algorithm which dynamically creates sampling regions using a distance metric in the CV space, thereby circumnavigating the computational expense resulting from the large number of bins in a high-dimensional CV space. Another approach termed resampling of ensembles by variation optimization (REVO)⁵² was proposed. REVO creates walker ensembles without regions by optimizing a metric that depends upon the pairwise distances between walkers.

Transition path sampling (TPS)⁴¹ is another technique which searches for dynamical bottlenecks faced by a system as it transitions between states. For instance, for a simple system with two states, A and B, there are many ways in which this transition can happen. Often, a random walk in the space of transition trajectories is used to create transition path ensembles. Variants of the same approach have been proposed such as milestoneing,⁵³ the string method,⁵⁴ etc.

Random resetting of trajectories, coined stochastic resetting,²⁴ has been recently applied to enhancing the sampling of MD simulations and has shown to result in an increase of an order of magnitude in long time scale processes for simpler systems. In stochastic resetting the set of short simulations is intermittently stopped and restarted where the resetting times are taken with constant steps, termed “sharp resetting”, or can be drawn from an exponential distribution, called “Poisson resetting”. The idea is that the resetting of the trajectories allows sampling of a richer set of pathways, which on average will lead to faster sampling.

It is pertinent to mention here that hybrid schemes can be adopted, as well. For instance in this foundational study⁵⁵ the authors used generalized ensemble methods, in particular simulated tempering,⁵⁶ to reach metastable states analyzed with Markov state models, and then an adaptive seeding

scheme was used to start new simulations from these broadly sampled states. This hybrid scheme allows one to cater for the rate limiting entropic barriers which are a significant issue in generalized ensemble methods where random walks in only the temperature space are performed. This also allows one to perform equilibrium sampling from nonequilibrium data, as the authors show that only local equilibrium is required, meaning that, even though individual trajectories only sample a subset of the metastable states, they cumulatively sample the complete space together.

RECENT ADVANCES

In this section, we focus on recent advances in adaptive sampling methodology as well as promising areas of future inquiry. In particular, we focus on deep-learning (DL)-based techniques. DL has revolutionized many scientific fields, but its impact in the molecular biosciences was catapulted by the high accuracy of recent protein structure prediction models.^{57–59} While not directly related to adaptive sampling, structure prediction models are useful to generate initial seeds for adaptive sampling algorithms⁶⁰ and are discussed in that context. Researchers have been concurrently working on ML models to analyze^{61–64} and accelerate^{65,66} MD simulations. Both sets of tools (structure prediction models and MD analysis/acceleration models) are becoming increasingly relevant to the field of adaptive sampling simulations. Here, we discuss how both types of models have been used to improve adaptive sampling MD simulations and how to connect both approaches.

AlphaFold v2 (AF2)⁵⁷ and adjacent models^{58,59} opened the gates to high-accuracy protein structure prediction from sequence data only. Although structural information is useful to researchers, the output from these models does not provide information on the thermodynamics or kinetics of the protein. For instance, a protein might exist in active and inactive conformational states that interconvert under equilibrium conditions, but AF2 might only predict one of the states. Moreover, we cannot make an inference about the relative free energies of the two states based on which one was predicted by AF2.⁶⁷ Any information about interconversion rates between states is also absent from these models. The lack of knowledge about the protein’s dynamics can make it impossible to infer its function or mechanism, and therefore complementary methods are required.

MD simulations are excellent tools to computationally resolve the dynamics of proteins, but as discussed in the [Introduction](#), the long-time-scale problem turns this approach impractical. As a way to mitigate this issue, researchers have used perturbation-based methods⁶⁸ to acquire diverse initial conformations for MD simulations from AF2 and similar models.^{60,69} Since in general the convergence of kinetic models is highly sensitive to the initial conditions,⁷⁰ adaptive sampling methods stand to gain great speedups simply by improving the prior knowledge of the system.

For brevity, we will restrict ourselves to two previous studies that applied perturbation-based methods on structure prediction models, but other recent works exist.⁷¹ The first study was used to gather initial structures for parallel simulations of *Plasmodium falciparum* plasmepsin II (PM II) with the intent to sample cryptic binding pockets.⁶⁰ The other method was used to obtain diverse structures of the Shwachman–Bodian–Diamond syndrome protein (SBDS) and the monocarboxylate transporter 1 (MCT1).⁶⁹ They

were not used for MD simulations, although they have the potential to be useful in this regard.

The first study perturbed the AF2 input by subsampling the multiple-sequence alignment (MSA) of the sequence of interest and enabling stochastic dropout. The process of MSA subsampling consists of restricting the depth of the alignment to a user-selected threshold and randomly selecting the sequence of clusters that will be represented. When passed through AF2, the stochastically subsampled MSAs will produce diverse predicted structures that capture some aspects of the dynamics implied by the experimental structures.⁶⁸ Enabling dropout eliminates a small percentage of the nodes in the neural network during a single forward pass, further perturbing the output of the model. The result of applying this technique to PM II was the prediction of a structural ensemble that partially sampled a known cryptic pocket in the protein. By launching parallel MD simulations from these structures and building a MSM from the trajectories, it was possible to recover the free energy landscape of pocket opening.⁶⁰ These simulations did not require any type of biasing force to find the same free energy basins that could be detected with biased methods.⁶⁰ This study shows that structure prediction models are useful to generate seeds for swarms of unbiased MD simulations. In terms of future directions, it might be interesting to “close the loop” and use the simulations to produce new restarting seeds through fine-tuning of the structure prediction model, as it has been shown that AF2 can be fine-tuned for different specific tasks.^{72,73}

A different method to perturb the input to structure prediction models consists of splitting the sequence into fragments, predicting the fragments' structures, and connecting them to form a protein. This resembles early approaches to structure prediction, such as Rosetta.⁷⁴ A more recent exploration of this approach was presented in MultiSFold,⁶⁹ a method that combines the variable-length fragment library (VFLib)⁷⁵ technique with multiple structure prediction models to diversify the predicted conformational ensemble. This method is particularly useful when different prediction models (e.g., AlphaFold and RosettaFold) predict different structures for the same protein. MultiSFold seems to be useful for interpolating structures between functional end points, as it was shown for SBDS and MCT1.⁶⁹ In the future, it might be interesting to use MultiSFold to find initial seeds for adaptive sampling and then construct new distograms from the simulations. These distograms could then be run through MultiSFold again to obtain new seeds and close the sampling cycle.

Besides using structure prediction models to obtain better initial seeds, DL offers other strengths that have proven useful for adaptive sampling simulations. In particular, the ability to approximate complex nonlinear functions from high-dimensional molecular trajectories makes DL models suitable for dimensionality reduction. Finding appropriate low-dimensional projections facilitates many tasks involved in adaptive sampling simulations (e.g., choice of reaction coordinates, clustering, etc.).

We will focus on three studies that have applied DL models to adaptive sampling,^{2,76,77} but of course related works exist.^{71,78,79} A connection between these three works is that all of them apply a similar simulation-training cycle that consists of launching the simulations, training the model on them, and then using the model to select seeds for new trajectories to restart the loop. Two of these works^{2,76} use

different types of variational autoencoders (VAEs)⁸⁰ as their base models, while the third one⁷⁷ uses primarily VAMP-Nets.⁶¹

Specifically, DeepDriveMD⁷⁶ employs a convolutional VAE (CVAE), while latent space-assisted adaptive sampling (LAST)² employs a VAE parametrized by a feed-forward network with fully connected layers. Both works use these deep learning (DL) models to learn low-dimensional latent representations that model the probability distribution of the structural ensemble. By selecting the outliers in such a distribution, one can recover the rare conformations that have not been sufficiently sampled, a goal similar to that of count-based adaptive sampling.²⁸ In LAST, a nonparametric kernel density estimate is applied to obtain the cumulative distribution function in the latent space; the lowest probability samples are then selected as the seeds for the next round of simulation. This is done differently in DeepDriveMD, where the data points are clustered in the latent space with the lowest reconstruction loss using DBSCAN⁸¹ and then the clusters with fewer than 10 members are selected to restart simulations (with a cap of 150 simulations maximum). The authors of DeepDriveMD also discuss practical aspects, such as distributing the simulations and ML training to different components in high-performance clusters.⁷⁶ LAST was shown to discover the conformational landscape of two proteins (adenyl cyclase⁸² and the VIVID flavoprotein⁸³) faster than structural dissimilarity sampling,⁸⁴ which is another adaptive sampling algorithm. It was estimated to take 40% of the time to explore the conformational landscape compared to long MD simulations, accounting for the training time of the VAE. Similarly, DeepDriveMD was tested on the Fs peptide⁸⁵ and it was found to provide a 2.33 folding speedup compared to parallel simulations.

MaxEnt VAMPNet⁷⁷ is the third approach to be discussed. In this technique, rather than identifying outliers in latent space, VAMPNet⁶¹ is used to classify the discovered conformations into metastable states. A VAMPNet is a deep learning model that encodes the entire mapping from molecular coordinates to Markovian states, thus reducing the traditional work process of transforming trajectory data into hand-crafted features and dimensionality reduction (Figure 5). Since VAMPNets can assign probabilities, p_i , of “belonging” to a metastable state to each conformation, the Shannon entropy (an information theoretic metric relating a distribution to the uncertainty of the model) is used to score the data points with the formula $H(\mathbf{p}) = -\sum_i p_i \log p_i$. The conformations with the highest entropy are selected for the next round of simulations. The rationale for selecting the structures based on their entropy is that the structures that cannot be clearly placed into a metastable state correspond to transition states or poorly sampled regions of the phase space. When applied to a small peptide (sequence WLALL⁸⁶), MaxEnt showed a 3× acceleration in conformational landscape discovery compared to a combination of VAMPNet and count-based adaptive sampling⁷⁷ and approximately 60% higher landscape discovery compared to reinforcement-learning-based approaches.^{47,48} Interestingly, the trajectories collected with MaxEnt also produced converged MSMs, while those collected with count-based sampling alone did not. When applied to a small protein, the villin headpiece,⁸⁷ it showed approximately 50% higher landscape discovery compared to the combination of VAMPNet and count-based sampling.

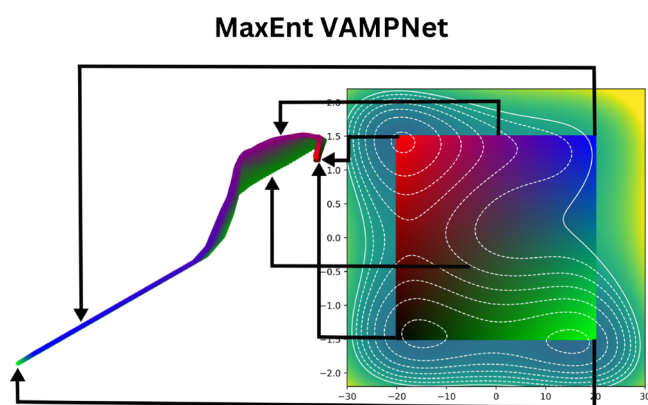


Figure 5. VAMPNets are deep learning models that learn nonlinear mapping to project the conformational landscape into a new space with kinetically relevant dimensions. The inset on the left depicts this latent space projection, where the colors on the internal rectangle drawn on the landscape reveal how this region is mapped by the model. In the technique termed MaxEnt VAMPNet, the inputs that maximize the Shannon entropy of the VAMPNet are chosen as seeds for restarting simulations.

In summary, this section covered how the DL can be used to improve adaptive sampling methods. Two avenues to incorporate DL into these workflows were explained: (1) using structure prediction models to augment prior knowledge of the simulated system and (2) training DL models to rank conformations for optimal seed selection. Combining these two approaches in different ways will probably yield new adaptive sampling methods that surpass the current state of the art. In the next section, we focus on the challenges that these approaches face.

■ CHALLENGES

Incorporating DL models into adaptive sampling algorithms already shows extremely encouraging results, but this also means that the MD workflows will inherit the challenges associated with these models. With respect to approaches such as DeepDriveMD, LAST, and MaxEnt VAMPNet, the main challenge of using deep models is that fitting them at each simulation round takes considerable computational time and power. Nonetheless, MD simulations continue to be slower, so even when accounting for the training time for DL models, there is a considerable speedup.² Another issue is that these models come accompanied by a host of design choices and optimization variables, although past results demonstrate speedups without an exhaustive hyperparameter tuning.^{2,48,76} Finally, the data sparsity toward the beginning of the simulation can result in noisy models whose validation scores might fluctuate widely. This could result in noisy selection criteria for simulation restarts. A potential way to prevent this issue would be to train several models initialized randomly and then combine their output using ensemble methods.⁸⁸

With respect to the use of structure prediction models to seed MD simulations, some limitations can also be noted. Assuming that the quality of the models is good enough to provide accurate conformations, there are issues when the predictions are too similar to those of a single native structure. When that is the case, launching simulations from several predicted structures might not provide an advantage, because the trajectories will be highly correlated. For instance, in the case of MultiSFold,⁶⁹ if RosettaFold⁵⁸ and AF2⁵⁷ output

similar structures, then the constructed optimization potentials will overlap and the generated ensemble will not be diverse. Similarly, for MSA subsampling of AF2, it might be challenging to gather diverse structures if few experimental structures are known.^{68,89} These issues might be tackled by future work through the combination of structure prediction models with MD simulation data sets to produce physically informed structure prediction of conformational ensembles.

We must also consider the possibility that the quality of the predicted model is not good enough to confidently employ it as an initial seed in MD simulations. If one starts MD simulations from unphysical conformations of a protein and then attempts to construct a MSM,¹⁵ this could result in disconnected states for which transitions cannot be sampled. One could establish a confidence threshold to accept or reject predicted structures prior to executing MD trajectories by using the predicted local distance difference test (pLDDT) score^{57,90} or similar metrics, but a benchmark for this specific purpose is lacking.

Overall, we have discussed some potential challenges in applying DL models in adaptive sampling algorithms. In particular, we noted that two of the most prominent challenges in the application of DL models in adaptive sampling algorithms are data sparsity and model validation. Potential avenues to mitigate these issues were mentioned.

■ CONCLUSIONS

In this Perspective, we have introduced the reader to many aspects of adaptive sampling MD simulations. Theoretically motivated approaches were presented in the [Theory](#) section. In the [Methods](#) section, several practiced approaches were summarized. Finally, we introduced recent advances that combine DL methods with adaptive sampling simulations in different ways and discussed their limitations and potential improvements.

While no quantitative comparisons between the methods were offered, it must be noted that there are still several points in which the community should agree before a rigorous comparison can be performed. First, an appropriate set of benchmark systems must be selected. To assess whether the tested methods work beyond toy examples, the complexity of the benchmarks should not be trivial. Moreover, the selection of the systems should obey some diversity criteria to ensure that the methods are robust. Second, the level of permissible initial knowledge about the systems must be set. Certain methods can exploit information about collective variables or known metastable states. It must be determined if we should judge the methods based on their ability to exploit known information or based on their “blind” performance. Third, the metrics used to rank the methods must be defined. A possible option would be the total simulation time taken to find the converged MFPT and free energy landscape. In this case, the definition of “converged” kinetic and thermodynamic metrics must be properly operationalized. Other factors come into play, as well. For methods involving training and inference from ML models, should this runtime be taken into account? Since most of the methods are parallel, should scaling and efficiency matter?

Although resolving these issues is outside the scope of this Perspective, the authors hope that the adaptive sampling community will agree to productive standards to propel the field forward. Finally, the authors hope that the ideas presented

in this work will inspire further innovation in the area of adaptive sampling for MD simulations.

AUTHOR INFORMATION

Corresponding Author

Diwakar Shukla – Center for Biophysics and Quantitative Biology, University of Illinois at Urbana–Champaign, Urbana, Illinois 61801, United States; Department of Bioengineering, Department of Chemical and Biomolecular Engineering, and Department of Plant Biology, University of Illinois at Urbana–Champaign, Urbana, Illinois 61801, United States; orcid.org/0000-0003-4079-5381; Email: diwakar@illinois.edu

Authors

Diego E. Kleiman – Center for Biophysics and Quantitative Biology, University of Illinois at Urbana–Champaign, Urbana, Illinois 61801, United States; orcid.org/0000-0002-3833-5872

Hassan Nadeem – Department of Bioengineering, University of Illinois at Urbana–Champaign, Urbana, Illinois 61801, United States

Complete contact information is available at:
<https://pubs.acs.org/10.1021/acs.jpcb.3c04843>

Notes

The authors declare no competing financial interest.

Biographies

Diego E. Kleiman is a PhD candidate in Biophysics and Quantitative Biology at the University of Illinois Urbana–Champaign supervised by Prof. Diwakar Shukla. He obtained his BS in Physics from New York University Abu Dhabi, UAE. His current research is focused on developing machine learning algorithms to accelerate adaptive sampling molecular dynamics simulations. He is also interested in investigating plant hormone receptors as targets for agrochemicals.

Hassan Nadeem is a Bioengineering PhD student in Prof. Diwakar Shukla's research group at the University of Illinois Urbana–Champaign. His research interests lie at the interface between molecular dynamics and data-driven computational approaches for biomolecular systems. Hassan obtained his BS in Mathematics from National University of Sciences & Technology, Pakistan, followed by an MS in Nuclear Engineering from Pakistan Institute of Engineering and Applied Sciences.

Diwakar Shukla is an associate professor in the Department of Chemical and Biomolecular Engineering at the University of Illinois at Urbana–Champaign. He is affiliated with the Center for Biophysics and Quantitative Biology, Department of Plant Biology and Bioengineering. His research work is focused on understanding complex biological processes using physics-based models and techniques. He received his BTech and MTech degrees in Chemical Engineering at the Indian Institute of Technology (IIT) Bombay, India. He then joined Massachusetts Institute of Technology where he received his MS and PhD degrees in Chemical Engineering for his work on solution biochemistry. Before joining University of Illinois, he worked with Prof. Vijay Pande at Stanford University on developing distributed computing approaches to understand protein dynamics.

ACKNOWLEDGMENTS

The authors acknowledge support from the National Science Foundation Early CAREER Award (NSF MCB-1845606).

D.E.K. was supported by a fellowship from The Molecular Sciences Software Institute under NSF grant CHE-2136142.

REFERENCES

- (1) MacKerell, A. D.; Bashford, D.; Bellott, M.; Dunbrack, R. L.; Evanseck, J. D.; Field, M. J.; Fischer, S.; Gao, J.; Guo, H.; Ha, S.; et al. All-Atom Empirical Potential for Molecular Modeling and Dynamics Studies of Proteins. *J. Phys. Chem. B* **1998**, *102*, 3586–3616.
- (2) Tian, H.; Jiang, X.; Xiao, S.; La Force, H.; Larson, E. C.; Tao, P. LAST: Latent Space-Assisted Adaptive Sampling for Protein Trajectories. *J. Chem. Inf. Model.* **2023**, *63*, 67–75.
- (3) Oostenbrink, C.; Villa, A.; Mark, A. E.; Gunsteren, W. F. V. A biomolecular force field based on the free enthalpy of hydration and solvation: The GROMOS force-field parameter sets 53A5 and 53A6. *J. Comput. Chem.* **2004**, *25*, 1656–1676.
- (4) Harder, E.; Damm, W.; Maple, J.; Wu, C.; Reboul, M.; Xiang, J. Y.; Wang, L.; Lupyan, D.; Dahlgren, M. K.; Knight, J. L.; et al. OPLS3: A Force Field Providing Broad Coverage of Drug-like Small Molecules and Proteins. *J. Chem. Theory Comput.* **2016**, *12*, 281–296.
- (5) Souza, P. C. T.; Alessandri, R.; Barnoud, J.; Thallmair, S.; Faustino, I.; Grünwald, F.; Patmanidis, I.; Abdizadeh, H.; Bruininks, B. M. H.; Wassenaar, T. A.; et al. Martini 3: a general purpose force field for coarse-grained molecular dynamics. *Nat. Methods* **2021**, *18*, 382–388.
- (6) Allen, M. P.; Tildesley, D. J. *Computer Simulation of Liquids*; Clarendon Press: Oxford, England, 1989.
- (7) Alessandri, R.; Souza, P. C. T.; Thallmair, S.; Melo, M. N.; de Vries, A. H.; Marrink, S. J. Pitfalls of the Martini Model. *J. Chem. Theory Comput.* **2019**, *15*, 5448–5460.
- (8) Hénin, J.; Lelièvre, T.; Shirts, M. R.; Valsson, O.; Delemotte, L. Enhanced Sampling Methods for Molecular Dynamics Simulations [Article v1.0]. *Living Journal of Computational Molecular Science* **2022**, *4*, 1583.
- (9) Donati, L.; Hartmann, C.; Keller, B. G. Girsanov reweighting for path ensembles and Markov state models. *J. Chem. Phys.* **2017**, *146*, No. 244112.
- (10) Wu, H.; Paul, F.; Wehmeyer, C.; Noé, F. Multiensemble Markov models of molecular thermodynamics and kinetics. *Proc. Natl. Acad. Sci. U. S. A.* **2016**, *113*, E3221–E3230.
- (11) Sugita, Y.; Okamoto, Y. Replica-exchange molecular dynamics method for protein folding. *Chem. Phys. Lett.* **1999**, *314*, 141–151.
- (12) Naritomi, Y.; Fuchigami, S. Slow dynamics in protein fluctuations revealed by time-structure based independent component analysis: The case of domain motions. *J. Chem. Phys.* **2011**, *134*, No. 065101.
- (13) Husic, B. E.; Pande, V. S. Markov state models: From an art to a science. *J. Am. Chem. Soc.* **2018**, *140*, 2386–2396.
- (14) Zuckerman, D. M.; Chong, L. T. Weighted ensemble simulation: review of methodology, applications, and software. *Annual review of biophysics* **2017**, *46*, 43–57.
- (15) Suárez, E.; Wiewiora, R. P.; Wehmeyer, C.; Noé, F.; Chodera, J. D.; Zuckerman, D. M. What Markov state models can and cannot do: Correlation versus path-based observables in protein-folding models. *J. Chem. Theory Comput.* **2021**, *17*, 3119–3133.
- (16) Aristoff, D.; Copperman, J.; Simpson, G.; Webber, R. J.; Zuckerman, D. M. Weighted ensemble: Recent mathematical developments. *J. Chem. Phys.* **2023**, *158*, No. 014108.
- (17) Baudel, M.; Guyader, A.; Lelièvre, T. On the Hill relation and the mean reaction time for metastable processes. *Stochastic Processes and their Applications* **2023**, *155*, 393–436.
- (18) Noé, F. Probability distributions of molecular observables computed from Markov models. *J. Chem. Phys.* **2008**, *128*, No. 244103.
- (19) Cao, S.; Montoya-Castillo, A.; Wang, W.; Markland, T. E.; Huang, X. On the advantages of exploiting memory in Markov state models for biomolecular dynamics. *J. Chem. Phys.* **2020**, *153*, No. 014105.
- (20) Dominic, A. J., III; Sayer, T.; Cao, S.; Markland, T. E.; Huang, X.; Montoya-Castillo, A. Building insightful, memory-enriched models

to capture long-time biochemical processes from short-time simulations. *Proc. Natl. Acad. Sci. U. S. A.* **2023**, *120*, No. e2221048120.

(21) Evans, M. R.; Majumdar, S. N. Diffusion with stochastic resetting. *Physical review letters* **2011**, *106*, No. 160601.

(22) Pal, A.; Reuveni, S. First passage under restart. *Physical review letters* **2017**, *118*, No. 030603.

(23) Reuveni, S. Optimal stochastic restart renders fluctuations in first passage times universal. *Physical review letters* **2016**, *116*, No. 170601.

(24) Blumer, O.; Reuveni, S.; Hirshberg, B. Stochastic Resetting for Enhanced Sampling. *Journal of physical chemistry letters* **2022**, *13*, 11230–11236.

(25) Starkov, D.; Belan, S. Universal performance bounds of restart. *arXiv preprint arXiv:2209.06611*, 2022.

(26) Shirts, M. R.; Pande, V. S. Mathematical analysis of coupled parallel simulations. *Physical review letters* **2001**, *86*, 4983.

(27) Hinrichs, N. S.; Pande, V. S. Calculation of the distribution of eigenvalues and eigenvectors in Markovian state models for molecular dynamics. *J. Chem. Phys.* **2007**, *126*, No. 244101.

(28) Weber, J. K.; Pande, V. S. Characterization and rapid sampling of protein folding Markov state model topologies. *J. Chem. Theory Comput.* **2011**, *7*, 3405–3411.

(29) Lane, T. J.; Shukla, D.; Beauchamp, K. A.; Pande, V. S. To milliseconds and beyond: challenges in the simulation of protein folding. *Curr. Opin. Struct. Biol.* **2013**, *23*, 58–65.

(30) Shukla, D.; Meng, Y.; Roux, B.; Pande, V. S. Activation pathway of Src kinase reveals intermediate states as targets for drug design. *Nat. Commun.* **2014**, *5*, 3397.

(31) Kohlhoff, K. J.; Shukla, D.; Lawrenz, M.; Bowman, G. R.; Konerding, D. E.; Belov, D.; Altman, R. B.; Pande, V. S. Cloud-based simulations on Google Exacycle reveal ligand modulation of GPCR activation pathways. *Nat. Chem.* **2014**, *6*, 15–21.

(32) Zimmerman, M. I.; Porter, J. R.; Ward, M. D.; Singh, S.; Vithani, N.; Meller, A.; Mallimadugula, U. L.; Kuhn, C. E.; Borowsky, J. H.; Wiewiora, R. P.; et al. SARS-CoV-2 simulations go exascale to predict dramatic spike opening and cryptic pockets across the proteome. *Nat. Chem.* **2021**, *13*, 651–659.

(33) Shukla, S.; Zhao, C.; Shukla, D. Dewetting Controls Plant Hormone Perception and Initiation of Drought Resistance Signaling. *Structure* **2019**, *27*, 692–702.e3.

(34) Chen, J.; White, A.; Nelson, D. C.; Shukla, D. Role of substrate recognition in modulating strigolactone receptor selectivity in witchweed. *J. Biol. Chem.* **2021**, *297*, No. 101092.

(35) Dutta, S.; Selvam, B.; Shukla, D. Distinct Binding Mechanisms for Allosteric Sodium Ion in Cannabinoid Receptors. *ACS Chem. Neurosci.* **2022**, *13*, 379–389.

(36) Chan, M. C.; Selvam, B.; Young, H. J.; Procko, E.; Shukla, D. The substrate import mechanism of the human serotonin transporter. *Biophys. J.* **2022**, *121*, 715–730.

(37) Feng, J.; Selvam, B.; Shukla, D. How do antiporters exchange substrates across the cell membrane? An atomic-level description of the complete exchange cycle in NarK. *Structure* **2021**, *29*, 922–933.e3.

(38) Selvam, B.; Yu, Y.-C.; Chen, L.-Q.; Shukla, D. Molecular Basis of the Glucose Transport Mechanism in Plants. *ACS Central Science* **2019**, *5*, 1085–1096.

(39) He, Z.; Paul, F.; Roux, B. A critical perspective on Markov state model treatments of protein-protein association using coarse-grained simulations. *J. Chem. Phys.* **2021**, *154*, No. 084101.

(40) Huber, G.; Kim, S. Weighted-ensemble Brownian dynamics simulations for protein association reactions. *Biophys. J.* **1996**, *70*, 97–110.

(41) Cerjan, C. J.; Miller, W. H. On finding transition states. *J. Chem. Phys.* **1981**, *75*, 2800–2806.

(42) Bowman, G. R.; Ensign, D. L.; Pande, V. S. Enhanced Modeling via Network Theory: Adaptive Sampling of Markov State Models. *J. Chem. Theory Comput.* **2010**, *6*, 787–794.

(43) Doerr, S.; Fabritiis, G. D. On-the-Fly Learning and Sampling of Ligand Binding by High-Throughput Molecular Simulations. *J. Chem. Theory Comput.* **2014**, *10*, 2064–2069.

(44) Shamsi, Z.; Moffett, A. S.; Shukla, D. Enhanced unbiased sampling of protein dynamics using evolutionary coupling information. *Sci. Rep.* **2017**, *7*, No. 12700.

(45) Zimmerman, M. I.; Bowman, G. R. FAST Conformational Searches by Balancing Exploration/Exploitation Trade-Offs. *J. Chem. Theory Comput.* **2015**, *11*, 5747–5757.

(46) Pérez, A.; Herrera-Nieto, P.; Doerr, S.; Fabritiis, G. D. AdaptiveBandit: A Multi-armed Bandit Framework for Adaptive Sampling in Molecular Simulations. *J. Chem. Theory Comput.* **2020**, *16*, 4685–4693.

(47) Shamsi, Z.; Cheng, K. J.; Shukla, D. Reinforcement learning based adaptive sampling: REAPing rewards by exploring protein conformational landscapes. *J. Phys. Chem. B* **2018**, *122*, 8386–8395.

(48) Kleiman, D. E.; Shukla, D. Multiagent reinforcement learning-based adaptive sampling for conformational dynamics of proteins. *J. Chem. Theory Comput.* **2022**, *18*, 5422–5434.

(49) Kahn, H.; Harris, T. E. Estimation of particle transmission by random sampling. *National Bureau of Standards Applied Mathematics Series* **1951**, *12*, 27–30.

(50) Dixon, T.; Uyar, A.; Ferguson-Miller, S.; Dickson, A. Membrane-mediated ligand unbinding of the PK-11195 ligand from TSPO. *Biophysical journal* **2021**, *120*, 158–167.

(51) Dickson, A.; Brooks, C. L. WExplore: Hierarchical Exploration of High-Dimensional Spaces Using the Weighted Ensemble Algorithm. *J. Phys. Chem. B* **2014**, *118*, 3532–3542.

(52) Donyapour, N.; Roussey, N. M.; Dickson, A. REVO: Resampling of ensembles by variation optimization. *J. Chem. Phys.* **2019**, *150*, 244112.

(53) Faradjian, A. K.; Elber, R. Computing time scales from reaction coordinates by milestoning. *J. Chem. Phys.* **2004**, *120*, 10880–10889.

(54) E, W.; Ren, W.; Vanden-Eijnden, E. String method for the study of rare events. *Phys. Rev. B* **2002**, *66*, No. 052301.

(55) Huang, X.; Bowman, G. R.; Bacallado, S.; Pande, V. S. Rapid equilibrium sampling initiated from nonequilibrium data. *Proc. Natl. Acad. Sci. U. S. A.* **2009**, *106*, 19765–19769.

(56) Marinari, E.; Parisi, G. Simulated Tempering: A New Monte Carlo Scheme. *Europhysics Letters (EPL)* **1992**, *19*, 451–458.

(57) Jumper, J.; Evans, R.; Pritzel, A.; Green, T.; Figurnov, M.; Ronneberger, O.; Tunyasuvunakool, K.; Bates, R.; Židek, A.; Potapenko, A.; et al. Highly accurate protein structure prediction with AlphaFold. *Nature* **2021**, *596*, 583–589.

(58) Baek, M.; DiMaio, F.; Anishchenko, I.; Dauparas, J.; Ovchinnikov, S.; Lee, G. R.; Wang, J.; Cong, Q.; Kinch, L. N.; Schaeffer, R. D.; et al. Accurate prediction of protein structures and interactions using a three-track neural network. *Science* **2021**, *373*, 871–876.

(59) Ahdriz, G.; Bouatta, N.; Kadyan, S.; Xia, Q.; Gerecke, W.; O'Donnell, T. J.; Berenberg, D.; Fisk, I.; Zanichelli, N.; Zhang, B.; et al. OpenFold: Retraining AlphaFold2 yields new insights into its learning mechanisms and capacity for generalization. *bioRxiv* **2022**, DOI: 10.1101/2022.11.20.517210.

(60) Meller, A.; Bhakat, S.; Solieva, S.; Bowman, G. R. Accelerating Cryptic Pocket Discovery Using AlphaFold. *J. Chem. Theory Comput.* **2023**, *19*, 4355–4363.

(61) Mardt, A.; Pasquali, L.; Wu, H.; Noé, F. VAMPnets for deep learning of molecular kinetics. *Nat. Commun.* **2018**, *9*, 5.

(62) Wehmeyer, C.; Noé, F. Time-lagged autoencoders: Deep learning of slow collective variables for molecular kinetics. *J. Chem. Phys.* **2018**, *148*, No. 241703.

(63) Sultan, M. M.; Pande, V. S. Automated design of collective variables using supervised machine learning. *J. Chem. Phys.* **2018**, *149*, No. 094106.

(64) McCarty, J.; Parrinello, M. A variational conformational dynamics approach to the selection of collective variables in metadynamics. *J. Chem. Phys.* **2017**, *147*, No. 204109.

- (65) Wang, D.; Wang, Y.; Chang, J.; Zhang, L.; Wang, H.; E, W. Efficient sampling of high-dimensional free energy landscapes using adaptive reinforced dynamics. *Nat. Comput. Sci.* **2022**, *2*, 20–29.
- (66) Guo, A. Z.; Sevgen, E.; Sidky, H.; Whitmer, J. K.; Hubbell, J. A.; de Pablo, J. J. Adaptive enhanced sampling by force-biasing using neural networks. *J. Chem. Phys.* **2018**, *148*, No. 134108.
- (67) Chakravarty, D.; Porter, L. L. AlphaFold2 fails to predict protein fold switching. *Protein Sci.* **2022**, *31*, No. e4353.
- (68) Del Alamo, D.; Sala, D.; Mchaourab, H. S.; Meiler, J. Sampling alternative conformational states of transporters and receptors with AlphaFold2. *Elife* **2022**, *11*, No. e75751.
- (69) Hou, M.; Jin, S.; Cui, X.; Peng, C.; Zhao, K.; Song, L.; Zhang, G. Protein multiple conformations prediction using multi-objective evolution algorithm. *bioRxiv* **2023**.
- (70) Bhatt, D.; Zhang, B. W.; Zuckerman, D. M. Steady-state simulations using weighted ensemble path sampling. *J. Chem. Phys.* **2010**, *133*, No. 014110.
- (71) Vani, B. P.; Aranganathan, A.; Wang, D.; Tiwary, P. AlphaFold2-RAVE: From Sequence to Boltzmann Ranking. *J. Chem. Theory Comput.* **2023**, *19*, 4351–4354.
- (72) Motmaen, A.; Dauparas, J.; Baek, M.; Abedi, M. H.; Baker, D.; Bradley, P. Peptide-binding specificity prediction using fine-tuned protein structure prediction networks. *Proc. Natl. Acad. Sci. U. S. A.* **2023**, *120*, No. e2216697120.
- (73) Bradley, P. Structure-based prediction of T cell receptor: peptide-MHC interactions. *eLife* **2023**, *12*, No. e82813.
- (74) Rohl, C. A.; Strauss, C. E.; Misura, K. M.; Baker, D. *Methods in enzymology*; Elsevier: 2004; Vol. 383, pp 66–93.
- (75) Feng, Q.; Hou, M.; Liu, J.; Zhao, K.; Zhang, G. Construct a variable-length fragment library for de novo protein structure prediction. *Briefings in Bioinformatics* **2022**, *23*, 23–bbac086.
- (76) Lee, H.; Turilli, M.; Jha, S.; Bhowmik, D.; Ma, H.; Ramanathan, A. Deepdrivemd: Deep-learning driven adaptive molecular simulations for protein folding. 2019 IEEE/ACM Third Workshop on Deep Learning on Supercomputers (DLS); 2019; pp 12–19.
- (77) Kleiman, D. E.; Shukla, D. Active Learning of the Conformational Ensemble of Proteins Using Maximum Entropy VAMPNets. *J. Chem. Theory Comput.* **2023**, *19*, 4377–4388.
- (78) Ribeiro, J. M. L.; Bravo, P.; Wang, Y.; Tiwary, P. Reweighted autoencoded variational Bayes for enhanced sampling (RAVE). *J. Chem. Phys.* **2018**, *149*, No. 072301.
- (79) Ojha, A. A.; Thakur, S.; Ahn, S.-H.; Amaro, R. E. DeepWEST: Deep learning of kinetic models with the Weighted Ensemble Simulation Toolkit for enhanced sampling. *J. Chem. Theory Comput.* **2023**, *19*, 1342–1359.
- (80) Kingma, D. P.; Welling, M. Auto-encoding variational bayes. *arXiv preprint arXiv:1312.6114*, 2013.
- (81) Ester, M.; Kriegel, H.-P.; Sander, J.; Xu, X. A density-based algorithm for discovering clusters in large spatial databases with noise. *KDD'96: Proceedings of the Second International Conference on Knowledge Discovery and Data Mining*. **1996**, 226–231.
- (82) Schlauderer, G.; Proba, K.; Schulz, G. Structure of a mutant adenylate kinase ligated with an ATP-analogue showing domain closure over ATP. *Journal of molecular biology* **1996**, *256*, 223–227.
- (83) Schwerdtfeger, C.; Linden, H. VIVID is a flavoprotein and serves as a fungal blue light photoreceptor for photoadaptation. *EMBO journal* **2003**, *22*, 4846–4855.
- (84) Harada, R.; Shigeta, Y. Efficient conformational search based on structural dissimilarity sampling: applications for reproducing structural transitions of proteins. *J. Chem. Theory Comput.* **2017**, *13*, 1411–1423.
- (85) McGibbon, R. T. Fs MD Trajectories; 2014; https://figshare.com/articles/dataset/Fs_MD_Trajectories/1030363/1 (accessed May 11, 2023).
- (86) Scherer, M. K.; Trendelkamp-Schroer, B.; Paul, F.; Pérez-Hernández, G.; Hoffmann, M.; Plattner, N.; Wehmeyer, C.; Prinz, J.-H.; Noé, F. PyEMMA 2: A software package for estimation, validation, and analysis of Markov models. *J. Chem. Theory Comput.* **2015**, *11*, 5525–5542.
- (87) Chiu, T. K.; Kubelka, J.; Herbst-Irmer, R.; Eaton, W. A.; Hofrichter, J.; Davies, D. R. High-resolution x-ray crystal structures of the villin headpiece subdomain, an ultrafast folding protein. *Proc. Natl. Acad. Sci. U. S. A.* **2005**, *102*, 7517–7522.
- (88) Jeffares, A.; Liu, T.; Crabbé, J.; van der Schaar, M. Joint Training of Deep Ensembles Fails Due to Learner Collusion. *arXiv preprint arXiv:2301.11323*, 2023.
- (89) Stein, R. A.; Mchaourab, H. S. SPEACH_AF: Sampling protein ensembles and conformational heterogeneity with AlphaFold2. *PLOS Computational Biology* **2022**, *18*, No. e1010483.
- (90) Mariani, V.; Biasini, M.; Barbato, A.; Schwede, T. IDDT: a local superposition-free score for comparing protein structures and models using distance difference tests. *Bioinformatics* **2013**, *29*, 2722–2728.