

Kinetics and the Folding of RNA and Other Polymers

by

Guangyao Zhou

B.Sc., Peking University; Beijing, China, 2012

M.Sc., Brown University; Providence, RI, 2014

A dissertation submitted in partial fulfillment of the
requirements for the degree of Doctor of Philosophy
in The Division of Applied Mathematics at Brown University

PROVIDENCE, RHODE ISLAND

May 2019

© Copyright 2019 by Guangyao Zhou

This dissertation by Guangyao Zhou is accepted in its present form
by The Division of Applied Mathematics as satisfying the
dissertation requirement for the degree of Doctor of Philosophy.

Date_____

Stuart Geman, Ph.D., Advisor

Recommended to the Graduate Council

Date_____

Charles Lawrence, Ph.D., Reader

Date_____

Govind Menon, Ph.D., Reader

Approved by the Graduate Council

Date_____

Andrew G. Campbell, Dean of the Graduate School

Vita

Education

- 2008–2012 **B.Sc. in Statistics and Probability, Peking University, China**
2009–2012 **B.A. in Economics (Double Major), Peking University, China**
2014 **M.Sc. in Applied Math, Brown University, USA**
2012–2018 **Ph.D. in Applied Math, Brown University, USA**

Experiences

- 2017.5–2018.8 **Applied Scientist Intern, Amazon Lab126, Cupertino, CA**

Publications

- **Guangyao Zhou**, M. Harrison, and S. Geman. Differentiable Inference and Compositionality in Probabilistic Generative Models. *In progress*, 2018b
- **Guangyao Zhou**, E. Borenstein, O. Livne, and A. Brandt. Multiscale Optimization for Stochastic Ill-conditioning in Deep Neural Networks. *In preparation*, 2018a

- Jackson Loper*, **Guangyao Zhou***, and S. Geman. Overcoming Entropic Barriers by Capacity Hopping. *Manuscript available upon request*, 2018
- **Guangyao Zhou**, J. Loper, Matthew Harrison, and S. Geman. Base-pair Ambiguity and the Kinetics of RNA Folding. *bioRxiv*, 2018c. doi: 10.1101/329698. URL <https://www.biorxiv.org/content/early/2018/05/25/329698>
- **Guangyao Zhou**, S. Geman, and J. M. Buhmann. Sparse Feature Selection by Information Theory. *IEEE International Symposium on Information Theory (ISIT)*, 2014, pages 926–930, 2014. doi: 10.1109/ISIT.2014.6874968
- **Guangyao Zhou**, Z. Lu, and Y. Peng. L1-graph Construction Using Structured Sparsity. *Neurocomputing*, 120:441–452, 2013

* indicates these authors contributed equally to the paper

Preface and Acknowledgments

First and foremost, I would like to thank my advisor Stuart Geman. He convinced me to be bold in taking the plunge into an area that I know nothing of, and guided me through building the work in this thesis from the ground up. This has been a profound and transformational experience, and it fundamentally changed my notion about what's possible in research. He taught me math, and so much more. He taught me how to think, how to learn, and in a sense, how to live. He has been a great inspiration, and a constant reminder for me to stay curious, to keep exploring, and to strive for something higher. He has been the key in the personal transformation I went through over the past six years, and I'm forever grateful for this experience.

I would like to thank Jackson Loper, for being there with me the whole time, and for building the thesis work together; Matt Harrison, for all the valuable suggestions for my research throughout the years; Joachim Buhmann and Achi Brandt, for giving me valuable research experiences outside my Ph.D. work; Chip Lawrence, for the helpful suggestions to my thesis work as an insider in computational biology; and Govind Menon, for serving on my committee and reading and discussing my thesis with me.

I would like to thank Anthony Ianiero, for helping me get out of my physical low, and for spending all those hours in the gym teaching me everything about bodybuilding and getting my fitness to a level I never perviously imagined.

I would like to thank those people who used to be in my life, for prompting me to change and to grow. I would like to thank all my friends at Brown, who keep me company through the years, and make me feel like part of a community.

I would like to thank my parents, for always being supportive, and for giving me the freedom to pursue my dream. And finally, my fiancée, Siyi Chen, for the enduring love and support, and for making this place home.

Abstract of “Kinetics and the Folding of RNA and Other Polymers”, by Guangyao Zhou, Ph.D., Brown University, May 2019

In this thesis, we study kinetics and the folding of RNAs and other polymers. In the first part, we study the problem from a statistical perspective. Thinking kinetically, we introduce quantitative measures of “ambiguity” and demonstrate their statistical relationships to native secondary structure and to qualitative distinctions between RNA families. In the second part, we focus on entropic barriers in molecular dynamics simulations. Entropic barriers, while ubiquitous and important in molecular dynamics, received relatively little attention in the literature, and their theoretical characterization and understanding are severely lacking. Here, using a simple toy model with a golf-course energy landscape and two targets, we look at entropic barriers from the perspective of hitting probabilities of the targets. We present rigorous theoretical results to establish global, approximately constant hitting probabilities as an essential feature of entropic barriers, and connect the global hitting probabilities to local information around the targets (in the form of capacities of local sets around the targets) to facilitate the understanding of entropic barriers. Inspired by the theoretical results, a method called Capacity Hopping (CHop for short), based on an efficient capacity estimation algorithm, is proposed for overcoming entropic barriers in molecular dynamics simulations. Extensive numerical experiments on a concrete 5-dimensional version of the toy model show that CHop is highly effective: it can be nearly as accurate as naive simulations in terms of estimating hitting probabilities, but 750 times faster.

Contents

Vita	iv
Preface and Acknowledgments	vi
1 Base-pair Ambiguity and the Kinetics of RNA Folding	1
1.1 Introduction	2
1.2 Results	6
1.2.1 Basic Notation and Definitions	6
1.2.2 A First Look at the Data: Shuffling Nucleotides	8
1.2.3 Formal Statistical Analyses	14
2 Overcoming Entropic Barriers by Capacity Hopping	22
2.1 Introduction	23
2.2 A Toy Model	27
2.3 Theory: ε -flatness and CHop Probabilities	30
2.3.1 Hitting Probability as a Global Property for Entropic Barriers	32
2.3.2 Approximating Global Hitting Probabilities Using Local Infor-	
mation	34
2.4 Application: CHop Method and Capacity Estimation	35
2.4.1 CHop Method	36
2.4.2 Efficient Estimation of Capacities	37
2.5 Numerical Experiments and Results	45
2.5.1 Sanity Checks on Flat Energy Landscape	46
2.5.2 Results on Nontrivial Energy Landscape	50
2.6 Conclusions and Future Directions	54
A Stationary Reversible Diffusion Processes	56
B Proof of the Main Theorem in Chapter 2	64
C Proof of the ε-flatness of the Toy Model in Chapter 2	70
D Proof of the lemma in Chapter 2	82

List of Tables

1.1	Comparative Secondary Structures: calibrated ambiguity indexes, by RNA family. The number of molecules, the median length (number of nucleotides), and the median α scores for T-S and D-S ambiguity indexes (Eqs 1.2.15 and 1.2.16) for each of the seven RNA families studied. RNA molecules from the first two families are active as single molecules (unbound); the remaining five are bound in ribonucleoproteins. Unbound RNA molecules have lower ambiguity indexes.	13
1.2	MFE Secondary Structures: calibrated ambiguity indexes, by RNA family. Identical to Table 1.1, except that the ambiguity indexes and their calibrations are calculated using the MFE secondary structures rather than those derived from comparison analyses. There is little evidence in the MFE secondary structures for lower ambiguity indexes among unbound RNA molecules.	14
1.3	Numbers of Positive Ambiguity Indexes, by family. #mol's: number of molecules; # $d_{T-S} > 0$ and # $d_{D-S} > 0$: numbers of positive T-S and D-S ambiguity indexes, secondary structures computed by <i>comparative analysis</i> ; # $d_{\tilde{T}-\tilde{S}} > 0$ and # $d_{\tilde{D}-\tilde{S}} > 0$: numbers of positive T-S and D-S ambiguity indexes, secondary structures computed by <i>minimum free energy</i>	19

List of Figures

1.1	Bound or Unbound? ROC performance of classifiers based on thresholding T-S and D-S ambiguity indexes. Small values are taken as evidence for molecules that are active as single entities (unbound), as opposed to parts of ribonucleoproteins (bound). Classifiers in the left two panels use comparative secondary structures to compute ambiguity indexes; those on the right use (approximate) minimum free energies. In each of the four experiments, a conditional p-value was also calculated, based only on the signs of the indexes and the null hypothesis that positive indexes are distributed randomly among molecules of all types as opposed to the alternative that positive indexes are more typically found among families of bound RNA. Under the null hypothesis, the test statistic is hypergeometric—see Eq 1.2.19. <i>Upper Left: $p = 6.1 \cdot 10^{-37}$; Lower Left: $p = 2.1 \cdot 10^{-9}$; Upper Right: $p = 0.0048$; Lower Right: $p = 0.76$.</i>	18
1.2	Comparative or MFE? As in Figure 1.1, each panel depicts the ROC performance of a classifier based on thresholding the T-S (top two panels) or D-S (bottom two panels) ambiguity indexes. Here, small values are taken as evidence for comparative as opposed to MFE secondary structure. Either index, T-S or D-S, can be used to construct a good classifier of the origin of a secondary structure for the unbound molecules in our data set (left two panels) but not for the bound molecules (right two panels). Conditional p-values were also calculated, using the hypergeometric distribution and based only on the signs of the indexes. In each case and the null hypothesis is that comparative secondary structures are as likely to lead to positive ambiguity indexes as are MFE structures, whereas the alternative is that positive ambiguity indexes are more typical when derived from MFE structures: <i>Upper Left: $p = 1.3 \cdot 10^{-13}$; Upper Right: $p = 0.078$; Lower Left: $p = 1.0 \cdot 10^{-6}$; Lower Right: $p = 0.026$.</i>	21
2.1	The sets and distances involved on the toy model.	28
2.2	We picked 100 random locations for $X_0 \notin \mathcal{B}(0, 1) \setminus (\tilde{A} \cup \tilde{B})$. For each location we used 2000 simulations to estimate the probability of hitting A before B . The histogram of the 100 estimates is shown below. The ε -flatness condition holds with $\varepsilon = 0.0425$, and the CHop probability gives good approximation to this narrow range of hitting probabilities.	48

2.3	We picked 100 random locations for $X_0 \notin \mathcal{B}(0, 1) \setminus (\tilde{A} \cup \tilde{B})$. For each location we used 2000 simulations to estimate the probability of hitting A before B . The histogram of the 100 estimates is shown below. The ε -flatness condition holds with $\varepsilon = 0.0470$, and the CHop probability gives good approximation to this narrow range of hitting probabilities.	52
-----	---	----

CHAPTER ONE

Base-pair Ambiguity and the Kinetics of RNA Folding

1.1 Introduction

Discoveries in recent decades have established a wide range of biological roles served by RNA molecules, in addition to their better-known role as carriers of the coded messages that direct ribosomes to construct specific proteins. Non-coding RNA molecules participate in gene regulation, DNA and RNA repair, splicing and self-splicing, catalysis, protein synthesis, and intracellular transportation [Morris and Mattick \(2014\)](#), [Kung et al. \(2013\)](#). To understand the mechanisms of these actions, emphasis has to be placed on the native structures, both secondary and tertiary. Our focus here will be on secondary structure, often a vital intermediary, and a useful abstraction for understanding the functions of non-coding RNA molecules [Higgs \(2000\)](#).

Because of the time-consuming nature of experimental determination of RNA structures, a considerable amount of work has been put into computational (as opposed to experimental) approaches. For secondary structure prediction, comparative analysis is the gold standard [Gutell et al. \(1992\)](#). But accurate comparative analysis typically requires a large number of homologous sequences, which may or may not be available. On the other hand, the prevailing *computational* approaches are based on an important and controversial assumption. Namely, that the structures of non-coding RNA molecules, *in vivo*, are in thermal equilibrium. If this were the case, and if the configuration energies could be accurately specified, then secondary (and in principle, even tertiary) structures could be explored through Monte Carlo sampling of the Gibbs equilibrium distribution. Alternatively, depending on the nature of the energy landscape, secondary structures could be approximated by identifying the minimum free energy (MFE) configuration [Mathews et al. \(1999\)](#), [Zuker et al. \(1999\)](#). However, even if we put aside the necessity for approximations, the biological

relevance of equilibrium configurations has been a source of misgivings at least since 1969, when Levinthal pointed out that the time required to equilibrate might be too long by many orders of magnitude [Levinthal \(1969\)](#). In light of these observations, and considering the “frustrated” nature of the folding landscape, many have argued that when it comes to structure prediction for macromolecules, kinetic accessibility is more relevant than equilibrium thermodynamics [Higgs \(2000\)](#), [Flamm and Hofacker \(2008\)](#), [Baker and Agard \(1994\)](#).

The primary structures of RNA molecules typically afford many opportunities to form short or medium-length stems¹, most of which do not participate in the native structure. This situation not only makes it hard for the biologist to accurately predict secondary structure, but equally challenges the molecule to avoid these “energetic traps.” Once formed, they require a large amount of energy (not to mention time) to be unformed. By what mechanisms might these traps be discouraged, if not actually avoided? Ideally, the only stems available in the primary structure would be those that participate in the secondary structure, but this is obviously too restrictive for all but the shortest or simplest of the structural RNA molecules. Still, there can be little doubt that the necessity for repeatable and efficient folding produces a selective pressure against the more disruptive ambiguities—e.g. a sequence that can form half of either of two stems, both of which are kinetically accessible and stable, but only one of which is native, or a sequence that is unpaired in the native structure but can nonetheless form a stem that is stereochemically inconsistent with the native structure. By this reasoning, which emphasizes kinetics rather than equilibrium, *per se*, we might expect to find information in the intra-molecular ambiguities about the folding process, or even the native structure.

We will introduce quantitative measures of ambiguity and demonstrate their

¹ By which we will mean sequences of G·U (“wobble pairs”) and/or Watson-Crick pairs.

statistical relationships to native secondary structure and to qualitative distinctions between RNA families. Using a somewhat arbitrary convention, we will refer to sequences of length four (four consecutive nucleotides) as *segments*. Keeping this convention in mind, each line of statistical evidence is based on the same formal definition of the “local ambiguity” of a given segment. This is simply the number of its complementary pairs in the molecule. When we refer to the *location* of a segment, we will mean the location of the first element of the segment, counting from 5′ to 3′, and when we refer to the local ambiguity of a location, we will mean the local ambiguity of the segment at that location. Local ambiguity is an intrinsic property of the primary structure. We will be interested in exposing its statistical relationships to any given candidate secondary structure of the same molecule. For this purpose, for any particular secondary structure we distinguish three different kinds of locations:

Single: Locations where all nucleotides in the corresponding segment are unpaired in the secondary structure;

Double: Locations where all nucleotides in the corresponding segment are paired in the secondary structure;

Transitional: Locations where some nucleotides in the corresponding segment are paired and others are unpaired in the secondary structure.

Double and *transitional* locations participate in the candidate secondary structure, while *single* locations do not. As a general rule, local ambiguities measured at any of these locations will increase with the length of the molecule—there are simply more opportunities to find complementary sequences. Instead of using the local ambiguities themselves, we focus on the *differences* between ambiguities in and around

stems (*double* and *transitional* locations) from those at unpaired (*single*) locations. In particular, we defined the “T-S ambiguity index” to be the difference between the average local ambiguity at *transitional* locations minus the average at *single* locations. We defined the “D-S ambiguity index” similarly, but adjusted for the evident bias at *double* locations, each of which, after all, has at least one complementary pair somewhere in the molecule. The D-S ambiguity index, then, is the average local ambiguity at *double* locations minus the average over those *single* locations that have at least one ambiguity.

Using only primary and (comparative-analysis) secondary structures, elementary counting statistics, and exact, distribution-free, tests, we will give evidence that both the T-S and D-S ambiguity indexes significantly separate two groups of non-coding RNA molecules: one group consists of RNA families that operate, *in vivo*, as single entities—the Group I and Group II Introns; the other group is made up of RNA families known to be active as protein-RNA complexes (i.e. as ribonucleoproteins)—the transfer-messenger RNAs (tmRNA), the RNAs of signal recognition particles (SRP RNA), the ribonuclease P family (RNase P), and the 16s and 23s ribosomal RNAs (16s and 23s rRNA). In particular, “unbound” RNA molecules, which perform their functions without being part of a larger nucleoprotein complex, have systematically lower T-S and D-S ambiguity indexes than the “bound” RNA molecules found in ribonucleoproteins. There are many possible explanations. Possibly, the unbound molecules are more sensitive to ambiguities and their energy traps, especially those ambiguities that involve the structurally critical *double* and *transitional* regions, than the bound molecules, whose secondary structures may well be influenced by their chemical relationships to proteins. Interestingly, this distinction between the ambiguity indexes of bound and unbound molecules largely disappears when MFE structures are used instead of comparative-analysis structures in defining *dou-*

ble, *transitional*, and *single* locations. In fact, for most unbound molecules we can classify a candidate secondary structure (was it derived from comparative analysis or by a minimum-free-energy calculation?) just by looking at the difference in the T-S or the D-S index under the two structures.

The *Results* section is organized as follows: we first develop some basic notation and definitions, and then present an exploratory and largely informal statistical analysis. This is followed by formal results comparing ambiguities in unbound versus bound molecules, and then by a comparison of the ambiguities implied by secondary structures derived from comparative analyses to those derived through minimization of free energy.

1.2 Results

1.2.1 Basic Notation and Definitions

Consider a non-coding RNA molecule with N nucleotides. Counting from 5' to 3', we denote the primary structure by

$$p = (p_1, p_2, \dots, p_N), \text{ where } p_i \in \{A, G, C, U\}, i = 1, \dots, N \quad (1.2.1)$$

and the secondary structure by

$$s = \{(j, k) : \text{nucleotides } j \text{ and } k \text{ are paired, } 1 \leq j < k \leq N\} \quad (1.2.2)$$

and we define the *segment at location i* to be

$$P_i = (P_{i,1}, P_{i,2}, P_{i,3}, P_{i,4}) = (p_i, p_{i+1}, p_{i+2}, p_{i+3}) \quad (1.2.3)$$

There is no particular reason for using segments of length four, and in fact all qualitative conclusions are identical with segment lengths three, four, or five, and, for that matter, whether or not we include the G·U wobble pairs. We chose to include them.

Which segments are viable complementary pairs to P_i ? The only constraint on location is that an RNA molecule cannot form a loop of two or fewer nucleotides. Let A_i be the set of all segments that are potential pairs of P_i :

$$A_i = \{P_j : 1 \leq j \leq i - 7 \text{ or } i + 7 \leq j \leq N - 3\} \quad (1.2.4)$$

We can now define the *local ambiguity function*,

$$a(p) = (a_1(p), \dots, a_{N-3}(p))$$

which is a vector-valued function of the primary structure p . The vector has one component, $a_i(p)$, for each segment P_i , which is, simply, the number of feasible segments that are complementary to P_i :

$$\begin{aligned} a_i(p) &= \#\{P \in A_i : P \text{ and } P_i \text{ are complementary}\} \\ &= \#\{P_j \in A_i : (P_{i,k}, P_{j,5-k}) \in \{(A, U), (U, A), (G, C), (C, G), (G, U), (U, G)\}, \\ &\quad k = 1, \dots, 4\} \end{aligned} \quad (1.2.5)$$

We want to explore the relationship between ambiguity and secondary structure.

We can do this conveniently, on a molecule-by-molecule basis, by introducing another vector-valued function, this time depending only on a purported secondary structure. Specifically, the new function assigns a descriptive label to each location (i.e. each nucleotide), determined by whether the segment at the given location is fully paired, partially paired, or fully unpaired.

Formally, given a secondary structure s , as defined in equation (2), and a location $i \in \{1, 2, \dots, N - 3\}$, let $f_i(s)$ be the number of nucleotides in P_i that are paired under s :

$$f_i(s) = \#\{j \in P_i : (j, k) \in s \text{ or } (k, j) \in s, \text{ for some } 1 \leq k \leq N\} \quad (1.2.6)$$

Evidently, $0 \leq f_i(s) \leq 4$. The “paired nucleotides function” is then the vector-valued function of secondary structure defined as $f(s) = (f_1(s), \dots, f_{N-3}(s))$. Finally, we use f to distinguish three types of locations (and hence three types of segments): location i will be labeled

$$\begin{cases} \textit{single} & \text{if } f_i(s) = 0 \\ \textit{double} & \text{if } f_i(s) = 4 \\ \textit{transitional} & \text{if } 0 < f_i(s) < 4 \end{cases} \quad i = 1, 2, \dots, N - 3 \quad (1.2.7)$$

1.2.2 A First Look at the Data: Shuffling Nucleotides

Our goal is to explore connections between ambiguities and basic characteristics of RNA families, as well as the changes in these relationships, if any, when using comparative, as opposed to MFE, secondary structures. For each molecule and each location i , the segment at i has been assigned a “local ambiguity” $a_i(p)$ that

depends only on the primary structure, and a label (*single*, *double*, or *transitional*) that depends only on the secondary structure. Since the local ambiguity, by itself, is strongly dependent on the length of the molecule, and possibly on other intrinsic properties, we proposed two *relative* ambiguity indexes: T-S and D-S, each of which depends on both the primary (p) and purported secondary (s) structures. Formally,

$$d_{\text{T-S}}(p, s) = \frac{\sum_{j=0}^{N-3} a_j(p) c_j^{\text{tran}}(s)}{\sum_{j=0}^{N-3} c_j^{\text{tran}}(s)} - \frac{\sum_{j=0}^{N-3} a_j(p) c_j^{\text{single}}(s)}{\sum_{j=0}^{N-3} c_j^{\text{single}}(s)} \quad (1.2.8)$$

where we have used c_i^{double} and c_i^{single} for indicating whether location i is *transitional* or *single* respectively. In other words, for each $i = 1, 2, \dots, N - 3$

$$c_i^{\text{tran}} = \begin{cases} 1, & \text{if location } i \text{ is } \textit{transitional} \\ 0, & \text{otherwise} \end{cases} \quad (1.2.9)$$

$$c_i^{\text{single}} = \begin{cases} 1, & \text{if location } i \text{ is } \textit{single} \\ 0, & \text{otherwise} \end{cases} \quad (1.2.10)$$

In short, the T-S ambiguity index is the difference in the averages of the local ambiguities at *transitional* sites and *single* sites. Since, as noted earlier, the local ambiguity at every *double* location, i , is at least one ($a_i \geq 1$), the D-S index involves only those *single* locations, j , for which it is also the case that $a_j \geq 1$:

$$d_{\text{D-S}}(p, s) = \frac{\sum_{j=0}^{N-3} a_j(p) c_j^{\text{double}}(s)}{\sum_{j=0}^{N-3} c_j^{\text{double}}(s)} - \frac{\sum_{j=0}^{N-3} a_j(p) \hat{c}_j^{\text{single}}(p, s)}{\sum_{j=0}^{N-3} \hat{c}_j^{\text{single}}(p, s)} \quad (1.2.11)$$

Here, $\hat{c}^{\text{single}}$, which evidently depends on both primary and secondary structure, indicates those *single* locations with at least one pair: for each $i = 1, 2, \dots, N - 3$

$$\hat{c}_i^{\text{single}} = \begin{cases} 1, & \text{if location } i \text{ is } \textit{single} \text{ and } a_i(p) \geq 1 \\ 0, & \text{otherwise} \end{cases} \quad (1.2.12)$$

Thinking kinetically, we might expect to find relatively small values of the T-S and D-S indexes ($d_{\text{T-S}}$ and $d_{\text{D-S}}$) in RNA families which fold and operate as individual entities—what we call here the unbound RNA. One reason for believing this is that larger numbers of partial matches for a given sequence in or around a stem would likely interfere with the *nucleation* of the native stem structure, potentially disrupting folding, and nucleation appears to be a critical and perhaps even rate-limiting step in folding. Indeed, the experimental literature [Pörschke \(1974a,b, 1977\)](#), [Mohan et al. \(2009\)](#) has long suggested that stem formation in RNA molecules is a two-step process. When forming a stem, there is usually a slow nucleation step, resulting in a few consecutive base pairs at a nucleation point, followed by a fast zipping step. Since the ambiguity indexes are functions of the secondary structure, s , our expectation are made, implicitly, under the assumption that s is correct. For the time being we will focus only on comparative structures, returning later to the questions about MFE structures raised in the *Introduction*.

How are we to gauge $d_{\text{T-S}}$ and $d_{\text{D-S}}$, and how are we to compare their values across different RNA families? Consider the following experiment: for a given RNA molecule we create a “surrogate” which has the same nucleotides, and in fact the same counts of *all* four-tuple segments as the original molecule, but is otherwise ordered randomly. If ACCU appeared eight times in the original molecule, then it appears eight times in the surrogate, and the same can be said of all sequences

of four successive nucleotides—the frequency of each of the 4^4 possible segments is preserved in the surrogate. If we also preserve the locations of the *transitional*, *double*, and *single* labels, then we can compute new values for $d_{\text{T-S}}$ and $d_{\text{D-S}}$, say $\tilde{d}_{\text{T-S}}$ and $\tilde{d}_{\text{D-S}}$, from the surrogate. If we produce many surrogate sequences then we will get a distribution of surrogate values $\tilde{d}$, one for each of the two ambiguity indexes, to which we can compare d . We made several experiments of this type—for each of the two indexes, T-S and D-S, within each of the seven RNA families (Group I and Group II Introns, tmRNA, SRP RNA, RNase P, and 16s and 23s rRNA).

To make this precise, consider an RNA molecule with primary structure p and comparative secondary structure s . Construct a segment “histogram function,” $\mathcal{H}(p)$, which is the vector of the number of times that each of the 4^4 possible segments appears in p .² Let $\mathcal{E}(p)$ be the set of all permutations of the ordering of nucleotides in p that preserve the frequencies of four-tuples:

$$\mathcal{E}(p) = \{p' : \mathcal{H}(p') = \mathcal{H}(p)\}$$

Clever algorithms exist for efficiently drawing independent samples from the uniform distribution on $\mathcal{E}$ —see [Kandel et al. \(1996\)](#), [Fitch \(1983\)](#), [Altschul and Erickson \(1985\)](#). In this paper, we used an efficient and flexible implementation of the Euler algorithm, called uShuffle [Jiang et al. \(2008\)](#). Let $p^{(1)}, \dots, p^{(K)}$ be K such samples, and let $d_{\text{T-S}}(p^{(1)}, s), \dots, d_{\text{T-S}}(p^{(K)}, s)$ and $d_{\text{D-S}}(p^{(1)}, s), \dots, d_{\text{D-S}}(p^{(K)}, s)$ be the corresponding T-S and D-S ambiguity indexes. The results are not vary sensitive to K , nor to the particular sample, provided that K is large enough. We used $K=10,000$. Finally, let $\alpha_{\text{T-S}}(p, s) \in [0, 1]$ and $\alpha_{\text{D-S}}(p, s) \in [0, 1]$ be the left-tail

² $\mathcal{H}$ is made a vector by fixing an arbitrary ordering of all possible segments.

empirical probabilities, under the distributions defined by the ensembles

$$d_{\text{T-S}}(p, s), d_{\text{T-S}}(p^{(1)}, s), \dots, d_{\text{T-S}}(p^{(K)}, s) \quad (1.2.13)$$

$$\text{and} \quad d_{\text{D-S}}(p, s), d_{\text{D-S}}(p^{(1)}, s), \dots, d_{\text{D-S}}(p^{(K)}, s) \quad (1.2.14)$$

of choosing an ambiguity index less than or equal to $d_{\text{T-S}}(p, s)$ and $d_{\text{T-S}}(p, s)$, respectively:

$$\alpha_{\text{T-S}}(p, s) = \frac{1 + \#\{k \in \{1, \dots, K\} : d_{\text{T-S}}(p^{(k)}, s) \leq d_{\text{T-S}}(p, s)\}}{1 + K} \quad (1.2.15)$$

$$\alpha_{\text{D-S}}(p, s) = \frac{1 + \#\{k \in \{1, \dots, K\} : d_{\text{D-S}}(p^{(k)}, s) \leq d_{\text{D-S}}(p, s)\}}{1 + K} \quad (1.2.16)$$

In essence, each α score is a self-calibrated ambiguity index.

It is tempting to interpret $\alpha_{\text{T-S}}(p, s)$ as a p-value from a conditional hypothesis test: Given s and $\mathcal{H}$, test the null hypothesis that $d_{\text{T-S}}(p, s)$ is statistically indistinguishable from $d_{\text{T-S}}(p', s)$, where p' is a random sample from $\mathcal{E}$. If the alternative hypothesis were that $d_{\text{T-S}}(p', s)$ is too small to be consistent with the null, then the null is rejected in favor of the alternative with probability $\alpha_{\text{T-S}}(p, s)$. The problem with this interpretation, and the analogous one for $\alpha_{\text{D-S}}(p, s)$, is that this null hypothesis violates the simple observation that given $\mathcal{H}$ there is information in s about p , whereas $p^{(1)}, \dots, p^{(K)}$ are independent of s given $\mathcal{H}$. A larger problem is that there is no reason to believe the alternative; we are more interested in *relative* than absolute ambiguity indexes. Thinking of $\alpha_{\text{T-S}}(p, s)$ and $\alpha_{\text{D-S}}(p, s)$ as a calibrated, intra-molecular indexes, we want to know how they vary across RNA families, and whether these variations depend on the differences between comparative and MFE structures.

Nevertheless, $\alpha_{\text{T-S}}(p, s)$ and $\alpha_{\text{D-S}}(p, s)$ are useful statistics for exploratory anal-

ysis. Table 1.1 provides summary data about the α scores for each of the seven RNA families. For each molecule in each family we use the primary structure and the comparative secondary structure, and $K=10,000$ samples from $\mathcal{E}$, to compute individual T-S and D-S α scores (Eqs 1.2.15 and 1.2.16). Keeping in mind that a smaller value of α represents a smaller *calibrated* value of the corresponding ambiguity index, $d(p, s)$, there is evidently a disparity between ambiguity indexes of RNA molecules that form ribonucleoproteins and those of RNA that are already active as individual molecules. As a group, the unbound molecules have systematically lower ambiguity indexes. As already noted, this observation is consistent with, and in fact anticipated by, the kinetic point of view. Shortly, we will support this observation with ROC curves and rigorous hypothesis tests.

family	number molecules	median length	median $\alpha_{\text{T-S}}$	median $\alpha_{\text{D-S}}$
Group I Introns	125	447	0.449	0.911
Group II Introns	38	951	0.178	0.695
tmRNA	404	363	0.925	0.987
SRP RNA	359	272	0.789	0.957
RNase P	414	328	0.923	0.982
16s rRNA	354	1526	0.983	1.000
23s rRNA	162	2896	1.000	1.000

Table 1.1: **Comparative Secondary Structures: calibrated ambiguity indexes, by RNA family.** The number of molecules, the median length (number of nucleotides), and the median α scores for T-S and D-S ambiguity indexes (Eqs 1.2.15 and 1.2.16) for each of the seven RNA families studied. RNA molecules from the first two families are active as single molecules (unbound); the remaining five are bound in ribonucleoproteins. Unbound RNA molecules have lower ambiguity indexes.

Does the MFE structure similarly separate single-entity RNA molecules from those that form ribonucleoproteins? A convenient way to explore this question is to recalculate and recalibrate the ambiguity indexes of each molecule in each of the seven families, but using the MFE in place of the comparative secondary structures. The results are summarized in Table 1.2. By comparison to the results shown from

Table 1.1, the separation of unbound from bound molecules nearly disappears under the MFE secondary structures. Possibly, the comparative structures, as opposed to the MFE structures, better anticipate the need to avoid energy traps in the folding landscape. Here too we will revisit the data using ROC curves and proper hypothesis tests.

family	number molecules	median length	median α_{T-S}	median α_{D-S}
Group I Introns	125	447	0.846	0.993
Group II Introns	38	951	0.827	0.998
tmRNA	404	363	0.866	0.98
SRP RNA	359	272	0.807	0.941
RNase P	414	328	0.955	0.998
16s rRNA	354	1526	0.994	1.000
23s rRNA	162	2896	1.000	1.000

Table 1.2: **MFE Secondary Structures: calibrated ambiguity indexes, by RNA family.** Identical to Table 1.1, except that the ambiguity indexes and their calibrations are calculated using the MFE secondary structures rather than those derived from comparison analyses. There is little evidence in the MFE secondary structures for lower ambiguity indexes among unbound RNA molecules.

1.2.3 Formal Statistical Analyses

The T-S and D-S ambiguity indexes ($d_{T-S}(p, s)$ and $d_{D-S}(p, s)$) are intra-molecular measures of the extent to which segments within single-stranded regions have more Watson-Crick and wobble pairings than segments that lie in and around subsequences that are double-stranded. As such, they depend on both p and any purported secondary structure, s . Based on calibrated versions of these indexes ($\alpha_{T-S}(p, s)$ and $\alpha_{D-S}(p, s)$) and employing the comparative secondary structure for s we found support for the idea that non-coding RNA molecules which are active as individual entities are more likely to have small ambiguity indexes than RNA molecules

destined to operate as part of ribonucleoproteins. Furthermore, the difference appears to be sensitive to the approach used for identifying secondary structure—there is little, if any, evidence in the MFE secondary structures for lower ambiguities among unbound molecules.

These qualitative observations can be used to formulate precise statistical hypothesis tests. Many tests come to mind, but perhaps the simplest and most transparent are based on nothing more than the molecule-by-molecule signs of the ambiguity indexes. Whereas ignoring the actual values of the indexes is inefficient in terms of information, and possibly also in the strict statistical sense, tests based on signs require very few assumptions and are, therefore, more robust to model misspecification. All of the p-values that we will report are based on the hypergeometric distribution, which arises as follows.

We are given a population of M molecules, each with a binary outcome measure $B_m \in \{-1, +1\}$, $m = 1, \dots, M$. There are two subpopulations of interest: the first M_1 molecules, making up population 1, and the next M_2 molecules, which make up population 2. (So $M_1 + M_2 = M$.) We observe n_1 plus values in population 1 and n_2 in population 2

$$n_1 = \#\{m \in \{1, 2, \dots, M_1\} : B_m = +1\} \quad (1.2.17)$$

$$n_2 = \#\{m \in \{M_1 + 1, M_1 + 2, \dots, M\} : B_m = +1\} \quad (1.2.18)$$

We suspect that population 1 has less than its share of plus ones, meaning that the $n_1 + n_2$ population of plus ones was not randomly distributed among the M molecules. To be precise, let N be the number of plus ones that appear from a draw, without replacement, of M_1 samples from $B_1, \dots, B_M$. Under the null hypothesis,

H_o , n_1 is a sample from the hypergeometric distribution on N :

$$\mathbb{P}\{N = n\} = \frac{\binom{M_1}{n} \binom{M_2}{n+1+n_2-n}}{\binom{M}{n_1+n_2}} \quad \max\{0, n_1 + n_2 - M_2\} \leq n \leq \min\{n_1 + n_2, M_1\} \quad (1.2.19)$$

The alternative hypothesis, H_a , is that n_1 is too small to be consistent with H_o , leading to a left-tail test with p-value $\mathbb{P}\{N \leq n_1\}$ (which can be computed directly or using a statistical package, e.g. *hygecdf* in MatLab).

It is by now well recognized that p-values should never be the end of the story. One reason is that *any* departure from the null hypothesis in the direction of the alternative, no matter how small, is doomed to be statistically significant, with arbitrarily small p-value, once the sample size is sufficiently large. In addition to reporting p-values, we will also display estimated ROC curves, summarizing performance of a related classification problem. There are eight such problems: Under each of the comparative and MFE secondary structures, how well can we use $d_{T-S}(p, s)$ or $d_{D-S}(p, s)$ to classify a given molecule as bound or unbound? And, given $d_{T-S}(p, s)$ or $d_{D-S}(p, s)$ for a molecule that is known to be bound or unbound, how well can we classify the secondary structure, s , as coming from a comparative analysis versus an MFE analysis?

1.2.3.1 Single-entity RNA Molecules versus Protein-RNA Complexes

Consider an RNA molecule, m , selected from one of the seven families in our data set, with primary structure p and secondary structure s , computed by comparative analysis. Given only the T-S ambiguity index of m (i.e. given only $d_{T-S}(p, s)$), how accurately could we classify m as being unbound (i.e. from the Group I or Group II Introns) as opposed to bound (i.e. from one of the five families tmRNA, SRP RNA,

RNase P, 16s rRNA or 23s rRNA)? The foregoing exploratory analysis suggests constructing a classifier that declares a molecule to be ‘unbound’ when $d_{\text{T-S}}(p, s)$ is small, e.g. $d_{\text{T-S}}(p, s) < t$, where the threshold t governs the familiar trade off between rates of “true positives” (an unbound m is declared ‘unbound’) and “false positives” (a bound m is declared ‘unbound’). Large values of t favor low rates of false positives at the price of low rates of true positives, whereas small values of t favor high rates of true-positives at the price of high rates of false positives. Since for each molecule m we have both the correct classification (unbound or bound) and the statistic d , we can estimate the ROC performance of our threshold classifier by plotting the empirical values of the pair, ($\#$ false positives, $\#$ true positives), for each value of t . The ROC curve for the two-category (unbound versus bound) classifier based on thresholding $d_{\text{T-S}}(p, s) < t$ is shown in the upper-left panel of Figure 1.1. Also shown is the estimated area under the curve (AUC=0.83), which has a convenient and intuitive interpretation, as it is equal to the probability that for two randomly selected molecules, m from the unbound population and m' from the bound population, the T-S ambiguity index of m will be smaller than the T-S ambiguity index of m' .

As mentioned earlier, we can also associate a traditional p-value to the problem of separating unbound from bound molecules, based on the T-S ambiguity indexes. We consider only the signs (positive or negative) of these indexes, and then test whether there are fewer than expected positive indexes among the unbound, as opposed to the bound, population. This amounts to computing $\mathbb{P}\{N \leq n_1\}$ from the hypergeometric distribution—Eq (1.2.19). The relevant statistics can be found in Table 1.3, under the column labels $\# \text{mol's}$ and $\#d_{\text{T-S}} > 0$. Specifically, $M_1 = 120 + 36 = 156$ (number of unbound molecules), $M_2 = 404 + 356 + 411 + 352 + 155 = 1678$ (number of bound molecules), $n_1 = 53 + 10 = 63$ (number of positive T-S indexes among

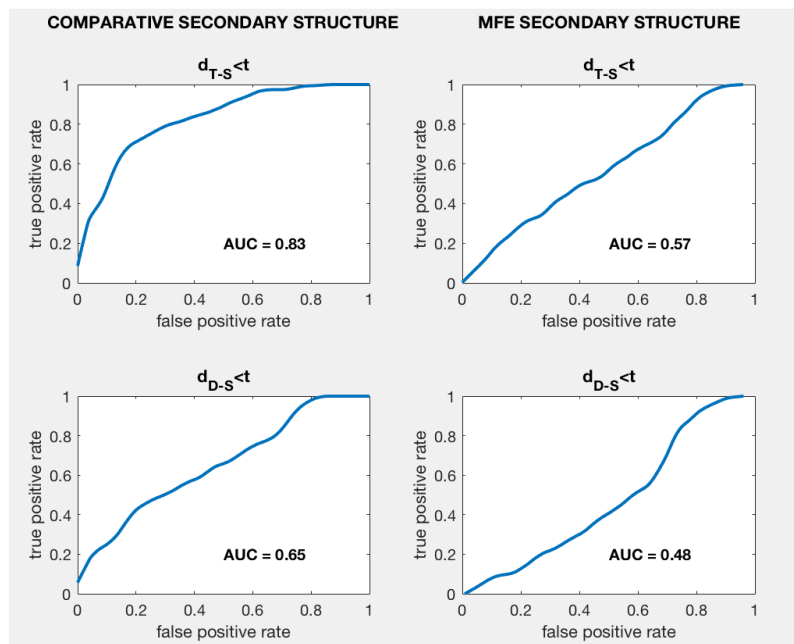


Figure 1.1: **Bound or Unbound?** ROC performance of classifiers based on thresholding T-S and D-S ambiguity indexes. Small values are taken as evidence for molecules that are active as single entities (unbound), as opposed to parts of ribonucleoproteins (bound). Classifiers in the left two panels use comparative secondary structures to compute ambiguity indexes; those on the right use (approximate) minimum free energies. In each of the four experiments, a conditional p-value was also calculated, based only on the signs of the indexes and the null hypothesis that positive indexes are distributed randomly among molecules of all types as opposed to the alternative that positive indexes are more typically found among families of bound RNA. Under the null hypothesis, the test statistic is hypergeometric—see Eq 1.2.19. *Upper Left:* $p = 6.1 \cdot 10^{-37}$; *Lower Left:* $p = 2.1 \cdot 10^{-9}$; *Upper Right:* $p = 0.0048$; *Lower Right:* $p = 0.76$.

unbound molecules) and $n_2 = 368 + 277 + 383 + 282 + 148 = 1458$. The resulting p-value, $6.1 \cdot 10^{-37}$, is essentially zero, meaning that the positive T-S indexes are not distributed proportional to the sizes of the unbound and bound populations, which is by now obvious, in any case. To repeat our caution, small p-values conflate sample size with effect size, and for that reason we have chosen additional ways, using permutations as well as classifications, to look at the data.

The comparative secondary structure of an RNA molecule, when combined with

family	#mol's	# $d_{T-S} > 0$	# $d_{D-S} > 0$	# $d_{\tilde{T}-\tilde{S}} > 0$	# $d_{\tilde{D}-\tilde{S}} > 0$
Group I Introns	120	53	103	97	117
Group II Introns	36	10	21	29	34
tmRNA	404	368	396	358	396
SRP RNA	356	277	319	273	306
RNase P	411	383	396	380	406
16s rRNA	352	282	323	326	349
23s rRNA	155	148	151	149	153

Table 1.3: **Numbers of Positive Ambiguity Indexes, by family.** #mol's: number of molecules; # $d_{T-S} > 0$ and # $d_{D-S} > 0$: numbers of positive T-S and D-S ambiguity indexes, secondary structures computed by *comparative analysis*; # $d_{\tilde{T}-\tilde{S}} > 0$ and # $d_{\tilde{D}-\tilde{S}} > 0$: numbers of positive T-S and D-S ambiguity indexes, secondary structures computed by *minimum free energy*.

its primary structure, can be used to construct a measure—the T-S ambiguity index—which distinguishes unbound from bound molecules with good accuracy. Can the same can be said for the D-S index? Yes, albeit with lower accuracy. To demonstrate, we followed the identical procedure, except that we assigned the index $d_{D-S}(p, s)$ rather than $d_{T-S}(p, s)$ to each molecule. The ROC curve (with area 0.65) is shown in the lower-left panel of Figure 1.1. The hypergeometric test, based on the thresholded (signed) values of d_{D-S} (positive counts, for each RNA family, can be found in Table 1.3) has p-value $2.1 \cdot 10^{-9}$. Finally, the right-hand panels in Figure 1.1 mirror the left-hand panels, except that all secondary structures were computed by (approximately) maximizing free energy rather than comparative analysis. The classification results are substantially less convincing and the p-values substantially higher.

1.2.3.2 Comparative Analysis versus Minimum Free Energy

As we have just seen, ambiguity indexes based on MFE secondary structures, as opposed to comparative secondary structures, do not make the same stark distinc-

tion between single and bound RNA molecules. To explore this a little further, we can turn the analyses of the previous paragraphs around and ask to what extent knowledge of the group of a molecule (unbound or bound) and its ambiguity index (e.g. the T-S ambiguity index) is sufficient to predict the source of the secondary structure—comparative or free energy?

Interestingly, there is a sharp difference in predicting the source of secondary structure in unbound molecules as opposed to bound molecules. Consider the top two ROC curves in Figure 1.2. In each of the two experiments a classifier was constructed by thresholding the T-S ambiguity index, declaring the secondary structure to be ‘comparative’ when $d_{\text{T-S}}(p, s) < t$ and ‘MFE’ otherwise, since, as we have seen, smaller values of the ambiguity indexes are generally associated with the comparative secondary structure. The ROC curves are then swept out by varying t . The difference between the two panels is in the population used for the classification experiments—unbound molecules in the upper-left and bound molecules in the upper-right. In unbound molecules, small values of the T-S ambiguity index are a good indication that a secondary structure was derived from comparative rather than free-energy analysis, but not in bound molecules. The corresponding hypothesis tests seek evidence against the null hypotheses that in a given group (unbound or bound) the set of positive T-S ambiguity indexes ($d_{\text{T-S}}(p, s) > 0$) are equally distributed between the comparative and free-energy derived indexes, and in favor of the alternatives that the T-S ambiguity indexes are less typically positive for the comparative secondary structures. The necessary data can be found in Table 1.3. The results are consistent with the classification experiments: the hypergeometric p-value is $1.3 \cdot 10^{-13}$ for the unbound population and 0.078 for the bound population.

The same experiments, with the same conclusions, were also performed using the D-S ambiguity index, as shown in the bottom two panels of Figure 1.2, for which

the corresponding hypergeometric p-values are $1.0 \cdot 10^{-6}$ (unbound population) and 0.026 (bound population).

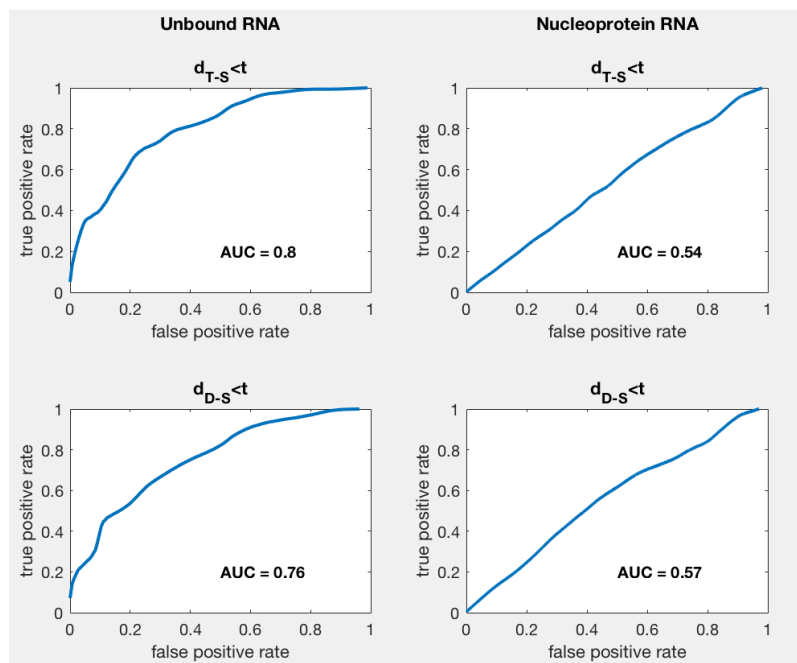


Figure 1.2: **Comparative or MFE?** As in Figure 1.1, each panel depicts the ROC performance of a classifier based on thresholding the T-S (top two panels) or D-S (bottom two panels) ambiguity indexes. Here, small values are taken as evidence for comparative as opposed to MFE secondary structure. Either index, T-S or D-S, can be used to construct a good classifier of the origin of a secondary structure for the unbound molecules in our data set (left two panels) but not for the bound molecules (right two panels). Conditional p-values were also calculated, using the hypergeometric distribution and based only on the signs of the indexes. In each case and the null hypothesis is that comparative secondary structures are as likely to lead to positive ambiguity indexes as are MFE structures, whereas the alternative is that positive ambiguity indexes are more typical when derived from MFE structures: *Upper Left*: $p = 1.3 \cdot 10^{-13}$; *Upper Right*: $p = 0.078$; *Lower Left*: $p = 1.0 \cdot 10^{-6}$; *Lower Right*: $p = 0.026$.

CHAPTER TWO

Overcoming Entropic Barriers by Capacity Hopping

2.1 Introduction

Molecular dynamics, first introduced in the 1970s [McCammon et al. \(1977\)](#), [Warshel and Levitt \(1976\)](#), has evolved into an indispensable tool in the study of structures and functions of macromolecules, and allows us to obtain a detailed understanding of the dynamics of a diverse range of biomolecular processes, including the folding of macromolecules into their native configurations [Scheraga et al. \(2007\)](#), and the conformational changes involved in their functioning [Hospital et al. \(2015\)](#).

Yet despite the substantial increase in speed and accuracy of molecular dynamics simulations over the years, we are still severely limited in the timescale we can access. Even with specialized hardwares, we can only achieve atomic-level simulations on timescales as long as milliseconds [Dror et al. \(2012\)](#), while the timescales for various biomolecular processes vary widely, and can last for seconds or longer [Naganathan and Muñoz \(2005\)](#), [Zemora and Waldsich \(2010\)](#). The challenges of molecular dynamics result from the complicated energy landscapes associated with different biomolecular systems, and can be summarized by two different kinds of barriers: the enthalpic barriers, which arise because of the difficulty of escaping local minima, and the entropic barriers, which arise because of the difficulty of reaching a small number of target configurations out of a vast number of possible configurations.

Great efforts have been put into extending the timescale accessible by molecular dynamics simulations. A lot of methods focus on enhanced sampling, with simulated annealing [Kirkpatrick et al. \(1983\)](#), genetic algorithms [Goldberg \(1989\)](#) and parallel tempering [Sugita and Okamoto \(1999\)](#) as representatives. However, dynamics is destroyed in these methods, and additional work [Yang et al. \(2007\)](#), [Andrec et al. \(2005\)](#), [Zheng et al. \(2009\)](#), [Huang et al. \(2010\)](#) is needed in order to recover the

kinetic information of the system, usually under the framework of kinetic transition networks [Noé et al. \(2006\)](#), [Wales \(2006\)](#) or Markov State Models [Pande et al. \(2010\)](#), [Chodera and Noé \(2014\)](#), [Husic and Pande \(2018\)](#). When focusing directly on kinetics, the foundational work is transition state theory [Eyring \(1935\)](#), [Chandler \(1978\)](#), [Wigner \(1997\)](#). Various other methods build on top of this, including transition path sampling [Dellago et al. \(1998\)](#), [Bolhuis et al. \(2002\)](#), transition interface sampling [van Erp and Bolhuis \(2005\)](#), and the more recent transition path theory [E. and Vanden-Eijnden \(2006\)](#), [E and Vanden-Eijnden \(2010\)](#). The essential idea here is to provide information about the reaction rates between different states, which can then be used under the framework of kinetic transition networks [Noé et al. \(2006\)](#), [Wales \(2006\)](#) or Markov State Models [Pande et al. \(2010\)](#), [Chodera and Noé \(2014\)](#), [Husic and Pande \(2018\)](#) to understand the kinetics of the system. In addition, different accelerated molecular dynamics methods [Perez et al. \(2009\)](#) (hyperdynamics [Voter \(1997\)](#), parallel replica [Voter \(1998\)](#) and temperature-accelerated dynamics [So/rensen and Voter \(2000\)](#)) aim at directly accelerating the molecular dynamics simulations, by either smoothing the energy landscape, exploring energy basins in parallel, or raising the temperature while maintaining the correct dynamics. We refer the readers to some other reviews [Klenin et al. \(2011\)](#), [Christen and van Gunsteren \(2008\)](#), [Perez et al. \(2009\)](#) for a more comprehensive picture of this area.

Careful inspection of the existing methods shows that a majority of the previous developments focus on enthalpic barriers, where the central picture consists of an energy landscape with a multitude of local minima separated by high energetic barriers. A less studied but arguably more important picture involves entropic barriers, which arise in an energy landscape with a golf-course potential, having a large flat region with several small targets in between, where an NP-complete problem can still be constructed [Baum \(1986\)](#), [Wille \(1987\)](#), even with essentially no local minima.

The importance of entropic barriers lie in their ubiquity in biomolecular systems, and their roles as the rate-limiting step for a variety of processes, as demonstrated by the strong experimental evidence [Teschner et al. \(1987\)](#), [Jacob et al. \(1999\)](#), [Goldberg and Baldwin \(1999\)](#), [Plaxco and Baker \(1998\)](#) that “folding rates are controlled by the rate of diffusion of pieces of the open coil in their search for favorable contacts” [McLeish \(2005\)](#). In fact, “the vast majority of the space covered by the energy landscape must actually be flat” [McLeish \(2005\)](#), indicating the presence and significance of entropic barriers.

The problem of entropic barriers has been identified in various places. In previous works [Baum \(1986\)](#), [Wille \(1987\)](#), [Machta \(2009\)](#), it’s pointed out that enhanced sampling methods such as simulated annealing and parallel tempering work poorly in the presence of entropic barriers. In accelerated molecular dynamics [So/rensen and Voter \(2000\)](#), the authors refer to this problem as the “low barrier problem”. In the review [Christen and van Gunsteren \(2008\)](#), the author takes a largely pessimistic view on this subject: “If these processes are intrinsically slow, i.e. require an extensive sampling of state space” (which is indeed the case with entropic barriers), “not much can be done to speed up their simulation without destroying the dynamics of the system.”

Some efforts have been made [Reguera and Rubí \(2001\)](#), [Lwin and Luo \(2005\)](#), [Mondal and Ray \(2010\)](#), [Huang et al. \(2009\)](#), [So/rensen and Voter \(2000\)](#) to deal with entropic barriers, yet theoretical characterization and understanding of entropic barriers are still severely lacking. The analysis of entropic barriers raises new questions. The typical question of interest for enthalpic barriers is: “How long will it take to cross the barrier?” By contrast, entropic barriers carry quite a different primary question: “Where will we end up next?” Or, more generally, “what is the probability of hitting various targets?” This motivates us to look at entropic barriers from the

perspective of hitting probabilities.

In this paper, we argue for a new point of view on entropic barriers: that the intrinsic slowness/long timescale involved in entropic barriers can be a blessing rather than a curse. Indeed, if the targets are small and the potential energy is relatively flat, the system would reach local equilibrium before it crosses the entropic barrier by hitting one of the targets. In this case, the starting point becomes irrelevant. The hitting probabilities remain approximately a constant, and become a global property of the entropic barriers. Our key insight is that we can characterize entropic barriers by establishing the hitting probabilities as a global property. Furthermore, we can understand this global property using only local information from the vicinities of the targets, which, it turns out, comes in the form of capacities. The contributions of this paper are twofold:

Theory We present rigorous theoretical results for the characterization of entropic barriers, by establishing global, approximately constant hitting probabilities as one of their essential features, and for the understanding of entropic barriers, by connecting their global hitting probabilities to capacities of local sets around the targets.

Application Inspired by our theoretical results, we present CHop as a general method to overcome entropic barriers, and further propose an efficient capacity estimation algorithm to address the central computational issue of CHop, before verifying the effectiveness of CHop through extensive numerical experiments.

The rest of the paper is organized as follows. In Section 2.2, we present a toy model with a golf-course potential that demonstrates the essential idea of an entropic

barrier. This example will run across the entire paper as a central thread for our theoretical and numerical analysis. In Section 2.3, we present our theoretical results. We propose the concept of ε -flatness to capture the idea that hitting probabilities remain approximately a constant, and establish the ε -flatness of the toy model in Theorem 2.3.2. We make further connections of the global hitting probabilities of entropic barriers to capacities of local sets around the targets in Corollary 1, for general stationary reversible SDEs. In Section 2.4, we talk about applications. We propose CHop as a general method to overcome entropic barriers, and talk about the central computational issue of CHop: the efficient estimation of the capacities. We present extensive numerical experiments and results in Section 2.5 to demonstrate the effectiveness of CHop, before giving conclusions and future directions in Section 2.6.

2.2 A Toy Model

To concretely demonstrate the CHop method, we present a simple toy model as an example of an entropic barrier. In this model, the potential energy landscape is flat, except for the regions around the two targets A and B . This large flat region forms an entropic barrier – the sheer size of it means that it would take the process a long time to find either A or B , even though there is no enthalpic barrier keeping it away from A or B . To further simplify the analysis and facilitate numerical validation, we will assume all sets of interest are spheres. A sketch of the relevant sets may be seen in Figure 2.1.

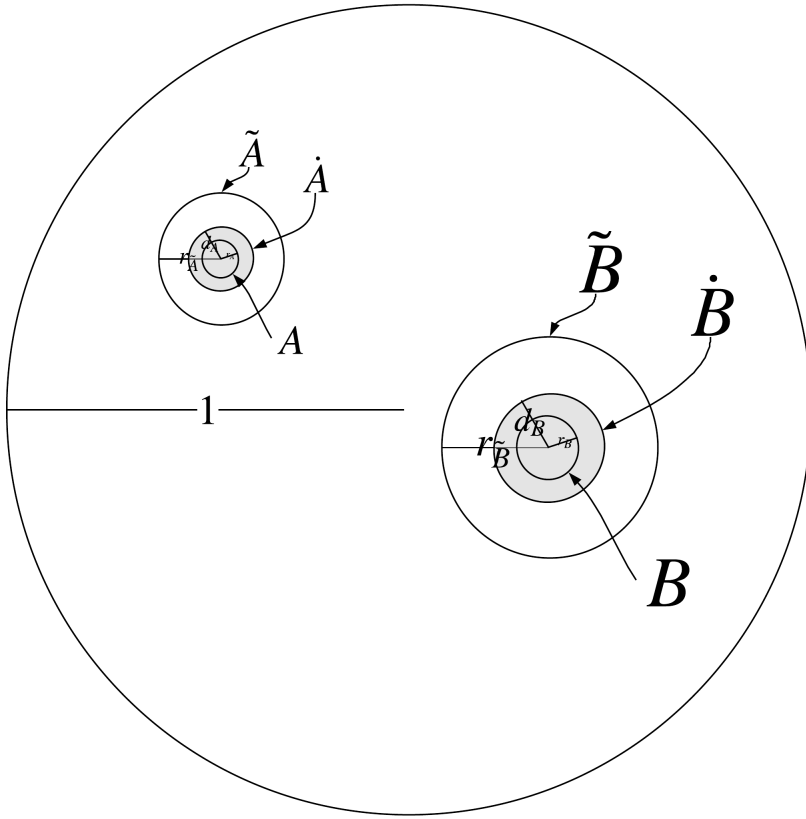


Figure 2.1: The sets and distances involved on the toy model.

Formally, denote the ball of radius r centered at a point $x \in \mathbb{R}^d$ by

$$\mathcal{B}(x, r) = \{y : \|y - x\| \leq r\}$$

Assume we have $x_A, x_B \in \Omega = \mathcal{B}(0, 1) \subset \mathbb{R}^d$ and $0 < r_A < d_A < r_{\tilde{A}}, 0 < r_B < d_B < r_{\tilde{B}}$ such that

$$r_{\tilde{A}} + \|x_A\| < 1, r_{\tilde{B}} + \|x_B\| < 1, \mathcal{B}(x_A, r_{\tilde{A}}) \cap \mathcal{B}(x_B, r_{\tilde{B}}) = \emptyset$$

Define

$$\begin{aligned} A &= \mathcal{B}(x_A, r_A), \dot{A} = \mathcal{B}(x_A, d_A), \tilde{A} = \mathcal{B}(x_A, r_{\tilde{A}}) \\ B &= \mathcal{B}(x_B, r_B), \dot{B} = \mathcal{B}(x_B, d_B), \tilde{B} = \mathcal{B}(x_B, r_{\tilde{B}}) \end{aligned}$$

Our toy model is governed by the SDE

$$dX_t = -\nabla U(X_t)dt + dW_t \tag{2.2.1}$$

with reflecting boundary at $\partial\mathcal{B}(0, 1)$, where U is a continuous potential function for which $U(x) = 0, \forall x \notin \dot{A} \cup \dot{B}$, and W_t is the d -dimensional Brownian motion. In other words, the toy model follows a diffusion process within $\mathcal{B}(0, 1)$ with trivial potential energy outside the set $\dot{A} \cup \dot{B}$, and trivial diffusion coefficients. Given an initial configuration, $X_0 = x$, we will be interested to know the probability that the diffusion hits $A = \mathcal{B}(x_A, r_A)$ before it hits $B = \mathcal{B}(x_B, r_B)$.

This toy model encapsulates the essential characteristics of an entropic barrier in the biomolecular system: if r_A, r_B are small, the trajectory of X will generally be spent well away from A or B , in the large “flat” region of the energy landscape that separates A from B . The nontrivial energy landscapes in the immediate vicinity

of A and B reflect the local nature of energetic interactions in the system under biological conditions, and the reflecting boundary $\partial\Omega$ captures the notion that not all configurations for the biomolecular system are sensible, because of, for example, the limits on bond lengths, angles and dihedral angles in the case of RNA molecules.

The method we propose in this paper is by no means limited to this example. However, this example is a simple one we can use to test our method, both analytically and numerically. Analytically, we are readily able to give some theorems for this special case. Numerically, even when r_A, r_B are very small and the dimension is modestly high, it is still (barely) possible to perform direct simulations (thanks in part to the very simple geometries involved). This enables us to test our method against the ground truth. Theorems and experiments in more general cases should be possible, but remain for future work.

2.3 Theory: ε -flatness and CHop Probabilities

Denote the m -dimensional Brownian motion by W_t . Assume the underlying physical process $\{X_t\}_{t \geq 0}$ that's governing the macromolecule is a stationary reversible diffusion process in the configuration space Ω with invariant measure $\mu = e^{-U(x)}dx$, and its dynamics can be described by the SDE

$$dX_t = b(X_t)dt + \sigma(X_t)dW_t \quad (2.3.1)$$

where $b(\cdot) : \Omega \rightarrow \mathbb{R}^d$ and $\sigma(\cdot) : \Omega \rightarrow \mathbb{R}^{d \times m}$ are differentiable vector-valued and matrix-valued functions. We refer the reader to Appendix A for a detailed account of stationary reversible diffusion processes. Fix some target sets A and B , our

primary interest will be in the theoretical understanding of entropic barriers from the perspective of hitting probabilities. Intuitively, an entropic barrier is a large (relative to the targets) subset of the configuration space Ω in which the energy $U(x)$ is relatively flat.

Let $\tau_S = \inf\{t \geq 0 : X_t \in S\}$, $\forall S \subset \mathbb{R}^d$ denote the time at which the process $\{X_t\}_{t \geq 0}$ first hits the set S . We care about hitting probabilities of the form

$$h_{A,B}(x) = \mathbb{P}(X_{\tau_{A \cup B}} \in A | X_0 = x), \forall x \in \Omega$$

which denotes the probability that the diffusion hits A before B given that it starts at configuration x . Our basic intuition is that, because of the large size (relative to the targets) of the entropic barrier, and the flatness of the energy landscape within the entropic barrier, the hitting probability is actually a global property of the entropic barrier: no matter where we start inside the entropic barrier, we would reach local equilibrium before we hit any of the targets, and the hitting probability remains approximately a constant. Furthermore, because we reach local equilibrium before we hit the targets, the internal structure of the entropic barrier doesn't matter, and we expect the global hitting probability can be accurately estimated using only local information around the targets.

In this section, we make the above intuitions rigorous: we introduce the concept of ε -flatness, to capture the notion that hitting probability remains approximately a constant, and is a global property for entropic barriers; we establish the connection between hitting probabilities and the capacities of sets around targets, to show that we can understand this global property of the entropic barrier using only local information around the targets.

2.3.1 Hitting Probability as a Global Property for Entropic Barriers

Definition 2.3.1. (ε -flatness) For any function u on Ω and an $\varepsilon > 0$, a set $M \subset \Omega$ is ε -flat for the function u if $\sup_{x,y \in M} |u(x) - u(y)| < \varepsilon$.

The concept of ε -flatness is introduced to capture the notion that the hitting probability remains approximately a constant. We will be particularly interested in the case where a large and general set is ε -flat for some hitting probability function. Our central observation is that, because of the flatness of the energy landscape (exactly flat in the case of our toy model) in the entropic barrier, and the small size of the targets relative to the entropic barrier, we can establish ε -flatness of the entropic barrier for the hitting probability function, which means the probability that we hit a particular target first is more or less constant, no matter where we start the diffusion within the set. This makes the hitting probability a global property within the entropic barrier.

It's beneficial to point out that not all ε -flat sets for hitting probability functions are entropic barriers. In particular, if we consider the isocommitor surfaces, i.e. surfaces on which the hitting probability is exactly a constant, we can see these sets are obviously ε -flat (with $\varepsilon = 0$), but they are in no sense entropic barriers. We are really only interested in those large and general sets for which we can still establish ε -flatness.

In what follows, we present a theorem in the context of a characteristic entropic barrier, our toy model, which establishes the ε -flatness of the set $\mathcal{B}(0, 1) / (\tilde{A} \cup \tilde{B})$ for the hitting probability function when the targets A, B are sufficiently small or

the dimension is high, for suitably picked sets $\tilde{A}, \tilde{B}$.

Theorem 2.3.2. *Let $u \triangleq h_{A,B}$, where the diffusion and A, B are defined by the toy model given in Equation 2.2.1 parametrized by $x_A, x_B, r_A, d_A, r_{\tilde{A}}, r_B, d_B, r_{\tilde{B}}$. For any fixed value of $d \geq 3$ and $r_{\tilde{A}}, r_{\tilde{B}}, \varepsilon > 0$, there exists a constant $c = c(d, r_{\tilde{A}}, r_{\tilde{B}}, \varepsilon)$ such that if $d_A, d_B < c$ then*

$$\sup_{x,y \in \mathcal{B}(0,1)/(\tilde{A} \cup \tilde{B})} |u(x) - u(y)| < \varepsilon$$

Likewise for any fixed value of $d_A, r_{\tilde{A}}, d_B, r_{\tilde{B}}, \varepsilon > 0$, there exists a constant $c = c(d_A, r_{\tilde{A}}, d_B, r_{\tilde{B}}, \varepsilon)$ such that if $d \geq c$ then the same holds.

We defer the proof of this theorem to Appendix C. Inspection of the proof demonstrates that the key for establishing ε -flatness is a proper separation of time scales: it takes a short time for the process to reach local equilibrium in the entropic barrier, because of the flatness of the energy landscape in the entropic barrier, while it takes a long time for the process to hit the targets starting in the entropic barrier, because of the size of the targets being small compared with the entropic barrier. In future work it would be helpful to find more general conditions under which ε -flatness can be established, but we expect that it would be possible to establish ε -flatness whenever we have reasonable flatness conditions for the energy landscape in a set that's large relative to the targets.

2.3.2 Approximating Global Hitting Probabilities Using Local Information

For a hitting probability function, if a set is ε -flat, then the hitting probability is a global property, and remains approximately a constant in this set. The key theoretical insight of this paper is that we can understand this global property using only local information around the targets. It turns out this local information comes in the form of capacities of sets around the targets. For $A \subset \tilde{A} \subset \Omega$, the capacity $\text{cap}(A, \tilde{A})$ under the stationary reversible diffusion process $\{X_t\}_{t \geq 0}$ is given by

$$\text{cap}(A, \tilde{A}) = \frac{1}{2} \int_{\Omega} \nabla h_{A, \tilde{A}^c}(x)^T a(x) \nabla h_{A, \tilde{A}^c}(x) e^{-U(x)} dx$$

where $a(x) = \sigma(x)\sigma(x)^T$ is the diffusion matrix. We refer the reader to Appendix A for more details on capacity and the related concept of Dirichlet form.

We start by considering the simple case where there are only two targets A and B , and we are interested to know which one we hit first. For the general stationary reversible SDE given by Equation 2.3.1, we have our main theorem

Theorem 2.3.3. *Assume we have $A, B \subset \Omega$ disjoint. Define*

$$u(x) = h_{A,B}(x) = \mathbb{P}(X_{\tau_{A \cup B}} \in A | X_0 = x)$$

$\forall \varepsilon \in (0, \frac{2}{9}]$, if we can find $\tilde{A}, \tilde{B} \subset \Omega$, s.t. $A \subset \tilde{A}, B \subset \tilde{B}$, and $M = \Omega / (\tilde{A} \cup \tilde{B})$ is ε -flat for the function u , then we have

$$\sup_{x \in M} \left| u(x) - \frac{\text{cap}(A, \tilde{A})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} \right| \leq \varepsilon + \sqrt{\frac{\varepsilon}{2}}$$

We defer the proof to Appendix B. This theorem establishes the connection of the global property of hitting probabilities to the local property of capacities. It generalizes naturally to the case of multiple targets:

Definition 2.3.4. Let $\{X_t\}_{t \geq 0}$ be a stationary reversible diffusion process on Ω . Pick $A_1 \cdots A_n \subset \Omega$ and $\tilde{A}_1 \cdots \tilde{A}_n \subset \Omega$ such that $A_k \subset \tilde{A}_k$. Then the **CHop Probability** for the set A_k (given all the other sets $A_1 \cdots, A_{k-1}, A_{k+1}, \cdots A_n \subset \Omega$ and $\tilde{A}_1 \cdots \tilde{A}_n \subset \Omega$) is given by

$$p_{A_k} = \frac{\text{cap}(A_k, \tilde{A}_k)}{\sum_{i=1}^n \text{cap}(A_i, \tilde{A}_i)}, k = 1, \dots, n$$

Corollary 1. Let $u_k = h_{A_k, \cup_i A_i}$ denote the probability the process hits A_k before the other target sets. Fix any $\varepsilon \in (0, \frac{2}{9}]$. If $M = \Omega \setminus \bigcup_{k=1}^n \tilde{A}_k$ is ε -flat for the functions $\{u_k\}_{k=1, \dots, n}$ then each u_k is well-approximated by the corresponding CHop Probability:

$$\sup_{x \in M} |u_k(x) - p_{A_k}| \leq \varepsilon + \sqrt{\frac{\varepsilon}{2}}, k = 1, \dots, n$$

This corollary follows immediately from Theorem 2.3.3 and additivity of the capacity (Proposition A.0.8 in Appendix A).

2.4 Application: CHop Method and Capacity Estimation

In this section, we focus on practical applications. Inspired by the theoretical insights above, we introduce the CHop method as a way to efficiently overcome entropic barriers in molecular dynamics simulations. We present the high-level idea of the

CHop method, and give a detailed account of the key computational issue of the CHop method, the efficient estimation of capacities. A general framework based on a simple lemma is given for capacity estimation, and the associated challenges are discussed. To demonstrate this framework, we give an outline of a flexible and general-purpose algorithm for capacity estimation, and make some concrete choices to show how exactly the algorithm works in the case of our toy model. This will serve as the basis for the numerical experiments and results in Section 2.5.

2.4.1 CHop Method

Recall that in the previous section on theory, we made the following two key points:

1. We can establish ε -flatness for entropic barriers, and the hitting probability remains approximately a constant, hence a global property of the entropic barrier.
2. We can accurately estimate the global hitting probabilities using CHop probabilities, which depend on capacities that can be determined solely based on the local information around the targets.

In most practical applications, molecular dynamics simulations of biomolecular systems are not fast enough to give us information on a biologically relevant time-scale. A large part of this difficulty results from entropic barriers [Teschner et al. \(1987\)](#), [Jacob et al. \(1999\)](#), [Goldberg and Baldwin \(1999\)](#), [Plaxco and Baker \(1998\)](#), [McLeish \(2005\)](#). The above theoretical insights provide us with an efficient way to overcome entropic barriers to get information of the biomolecular system on a much longer time-scale: we make use of our knowledge of the targets (which is typically

available, e.g. in the form of different secondary structures for RNAs) to carry out local simulations around the targets for the estimation of capacities. Capacities can then be used to calculate the CHop probabilities, which give us an accurate estimate of the global hitting probabilities. This enables us to make probabilistic jumps to one of the targets to accelerate the molecular dynamics simulations, which is the essence of the CHop method.

To summarize, on a high level, the CHop method consists of the following steps:

1. Identify an entropic barrier and establish its ε -flatness
2. Estimate the capacities of the sets around the targets to get the CHop probabilities
3. Make probabilistic jumps from within the entropic barrier to one of the targets, based on the approximated hitting probabilities given by the CHop probabilities

It's easy to see that, to make the CHop method practical, we need to be able to estimate the capacities accurately and efficiently. This would be our topic for the rest of this section.

2.4.2 Efficient Estimation of Capacities

We start our discussion on estimating the capacities with a simple lemma:

Lemma 2.4.1. *For a stationary reversible diffusion process X_t in Ω with invariant measure $\mu = e^{-U(x)}dx$ and diffusion matrix $a(x)$, given non-empty sets $A \subset G \subset$*

$\tilde{G} \subset \tilde{A}$ with smooth boundaries, the capacity $\text{cap}(A, \tilde{A})$ can be calculated by

$$\begin{aligned} \text{cap}(A, \tilde{A}) &= \int_{\partial \tilde{G}} h_{A, \tilde{A}^c}(x) e^{-U(x)} [a(x) \nabla h_{G, \tilde{G}^c}(x)]^T \mathbf{n}(x) dS \\ &\quad - \int_{\partial G} h_{A, \tilde{A}^c}(x) e^{-U(x)} [a(x) \nabla h_{G, \tilde{G}^c}(x)]^T \mathbf{n}(x) dS \end{aligned}$$

where $\mathbf{n}(x)$ in the integral is taken as the outward-facing normal vector at point x on the surface we are integrating on, and dS represents the surface integral.

We defer the proof of this lemma to Appendix D. This lemma provides us with a general framework for estimating the capacities. In particular, we can approximate the integrals in the lemma with a Monte Carlo method, as long as we can overcome three challenges:

1. To sample $y_1, y_2, \dots, y_m$ on ∂G and $\tilde{y}_1, \tilde{y}_2, \dots, \tilde{y}_n$ on $\partial \tilde{G}$ according to the invariant measure $e^{-U(x)} dx$ restricted on $\partial G, \partial \tilde{G}$ and compute the surface areas $|\partial G|$ and $|\partial \tilde{G}|$
2. To estimate $\nabla h_{G, \tilde{G}^c}$ on ∂G and $\partial \tilde{G}$
3. To estimate $h_{A, \tilde{A}^c}$ on ∂G and $\partial \tilde{G}$

With these pieces in place, using the lemma, we can approximate the capacity by

$$\begin{aligned} \text{cap}(A, \tilde{A}) \approx & |\partial \tilde{G}| \frac{\sum_i^n h_{A, \tilde{A}^c}(\tilde{y}_i) [a(\tilde{y}_i) \nabla h_{G, \tilde{G}^c}(\tilde{y}_i)]^T \mathbf{n}(\tilde{y}_i)}{\sum_i^n e^{U(\tilde{y}_i)}} \\ & - |\partial G| \frac{\sum_i^m h_{A, \tilde{A}^c}(y_i) [a(y_i) \nabla h_{G, \tilde{G}^c}(y_i)]^T \mathbf{n}(y_i)}{\sum_i^n e^{U(y_i)}} \end{aligned} \quad (2.4.1)$$

2.4.2.1 Overcoming the Challenges

The first challenge (sampling of points) can be addressed using existing methods, e.g. an importance sampling based approach, and is further simplified by the fact that we have the freedom to pick G and $\tilde{G}$. In some special cases (e.g. the toy model), it's not hard to pick $G, \tilde{G}$ in such a way that we can get exact samples from the invariant measure $e^{-U(x)}dx$.

The second challenge (estimation of $h_{G,\tilde{G}}$) can be addressed by making judicious choice of G and $\tilde{G}$ so that $\nabla h_{G,\tilde{G}^c}$ is relatively straightforward to compute. For example, in the general case, we can pick $\tilde{G}$ to be a small dilation of G (i.e. $\tilde{G} = \{x : \exists y \in G : |x - y| < \varepsilon\}$), so that the gradient can be well approximated by the surface normal of ∂G and $\partial \tilde{G}$. In some special cases (e.g. the toy model), we can even get analytical formulas for $\nabla h_{G,\tilde{G}^c}$.

The third challenge (estimation of $h_{A,\tilde{A}^c}$) requires more of a customized solution. The probabilistic interpretation of $h_{A,\tilde{A}^c}$ allows us to estimate it with local simulations. Since $\tilde{A}$ is a small, local space, $\tau_{A \cup \tilde{A}}$ will be relatively small, and we can successfully complete these simulations. In theory it would be possible to use many such simulations to estimate $h_{A,\tilde{A}^c}$. However, there are two main issues which tend to make this approach infeasible:

1. *Extreme probabilities.* If $h_{A,\tilde{A}^c}(x)$ is close to 0 or 1, it becomes much harder to obtain accurate estimates, since we would have to run a large number of trial simulations in order to get a good estimate of the hitting probability.
2. *Many starting points required.* The Monte Carlo approach requires us to run simulations for a collection of starting points on a particular surface, which in

turn requires us to run a huge number of trial simulations.

In this paper, we propose a flexible and general-purpose method for the efficient estimation of $h_{A,\tilde{A}^c}(x)$ on a surface ∂S , where $A \subset S \subset \tilde{A}$. The basic idea is to approximate the continuous diffusion with a discrete Markov process, and make use of the corresponding embedded Markov chain to estimate the hitting probability. In this approach, we only need to estimate the transition probabilities between different discretized states. The local transition probabilities of the Markov process tend to be much less extreme, solving issues of extreme probabilities. This approach also allows us to simultaneously estimate the hitting probabilities of all the discretized states on ∂S , which makes it computationally tractable to deal with all of the starting points that we need.

The method is closely related to milestoning [West et al. \(2007\)](#), [Bello-Rivas and Elber \(2015\)](#), [Aristoff et al. \(2016\)](#) and Markov state models [Pande et al. \(2010\)](#), [Chodera and Noé \(2014\)](#), [Husic and Pande \(2018\)](#), yet is more specialized and tailored to the problem at hand. In particular, we seek not to approximate the underlying process, but only to understand the hitting probability. Our goal is to make the method adaptive to the energy landscape, without the need for any prior knowledge. To this end, we evolve an ensemble of samples, and use a clustering-based approach to define the states.

The method has three main stages: determining the discretized states, running local simulations to estimate the transition probabilities, and estimating the hitting probabilities. We detail the method below, as part of the outline for our capacity estimation algorithm.

2.4.2.2 Outline of the Capacity Estimation Algorithm

In what follows, we put all the pieces from our previous discussions together, and present the outline of a flexible and general purpose capacity estimation algorithm:

Algorithm 2.4.2. *Estimating $\text{cap}(A, \tilde{A})$ for $A \subset \tilde{A} \subset \Omega$*

Input $A \subset \tilde{A} \subset \Omega$ and a stationary reversible diffusion process $\{X_t\}_{t \geq 0}$ in Ω with invariant measure $\mu = e^{-U(x)}dx$ and diffusion matrix $a(x)$

Output An estimated value of $\text{cap}(A, \tilde{A})$

1. Pick $G, \tilde{G}$ with smooth boundaries, s.t. $A \subset G \subset \tilde{G} \subset \tilde{A}$
2. Estimate the surface areas $|\partial G|$ and $|\partial \tilde{G}|$
3. Sample $y_1, y_2, \dots, y_m$ on ∂G and $\tilde{y}_1, \tilde{y}_2, \dots, \tilde{y}_n$ on $\partial \tilde{G}$ according to the invariant measure $e^{-U(x)}dx$ restricted on $\partial G, \partial \tilde{G}$, and estimate $h_{A, \tilde{A}^c}(y_i), i = 1, \dots, m$ and $h_{A, \tilde{A}^c}(\tilde{y}_i), i = 1, \dots, n$ using Algorithm 2.4.3 applied to the cases $S = G$ and $S = \tilde{G}$
4. Estimate $\nabla h_{G, \tilde{G}^c}(y_i), \mathbf{n}(y_i), i = 1, \dots, m$ and $\nabla h_{G, \tilde{G}^c}(\tilde{y}_i), \mathbf{n}(\tilde{y}_i), i = 1, \dots, n$, where $\mathbf{n}(x)$ denotes the outward-facing normal vector at point x on the corresponding surfaces
5. Estimate $\text{cap}(A, \tilde{A})$ using the formula given in Equation 2.4.1

where the algorithm for estimating local hitting probabilities is outlined as follows:

Algorithm 2.4.3. *Estimating $h_{A, \tilde{A}^c}(x)$ for many values of x along a shell ∂S*

Input $A \subset S \subset \tilde{A} \subset \Omega$ and a stationary reversible diffusion process $\{X_t\}_{t \geq 0}$ in Ω with invariant measure $\mu = e^{-U(x)} dx$. We also require a series of subsets

$$\tilde{A} \supset S_0 \supset S_1 \supset \cdots \supset S_{m-1} \supset S_m = S \supset S_{m+1} \supset \cdots \supset S_n = A$$

which indicate a kind of reaction coordinate.

Output A collection of points $z_1, \dots, z_{N_p}$ on ∂S sampled from the invariant measure $\mu = e^{-U(x)} dx$ restricted on ∂S , along with estimates of $h_{A, \tilde{A}^c}(y_i)$ for each point.

1. Discretize the space.

- (a) Generate an ensemble of samples $z_1, \dots, z_{N_p}$ on ∂S according to the invariant measure $\mu = e^{-U(x)} dx$ restricted to ∂S .
- (b) Evolve the ensemble on ∂S , by repeatedly taking a random sample from $\{z_1, \dots, z_{N_p}\}$ and follow the dynamics of $\{X_t\}_{t \geq 0}$ until the trajectory hits either ∂S_{m-1} or ∂S_{m+1} . Record the hitting locations on ∂S_{m-1} and ∂S_{m+1} until we have N_p samples on both ∂S_{m-1} and ∂S_{m+1} . Store the redundant trials for future estimation of transition probabilities.
- (c) Sequentially evolve the ensembles on $\partial S_{m-1}, \dots, \partial S_2$ and on $\partial S_{m+1}, \dots, \partial S_{n-2}$, until we have N_p samples on all of the intermediate surfaces $\partial S_1, \dots, \partial S_{n-1}$. Store the redundant trials for future estimation of transition probabilities.
- (d) For each one of the surfaces $\partial S_1, \dots, \partial S_{n-1}$, cluster the N_p samples on that surface into N_b states.

The result of this step is a partitioning of each shell ∂S_i into N_b discrete regions.

2. Estimate the transition probabilities between these states by running an additional N_s local simulations for each one of the N_b states on each surface. The

result of this step is an estimate of the probability of transitioning from state k on ∂S_i to state l on ∂S_j , which we denote by $P_{k,l}^{(i,j)}$, where $k, l \in \{1, \dots, N_b\}$ and $i, j \in \{1, \dots, n-1\}$ with $|i-j| = 1$.

3. Use the transition probabilities to get an estimate of the hitting probabilities for the N_b states on ∂S . In line with related works on Markov state models, [Pande et al. \(2010\)](#), [Chodera and Noé \(2014\)](#), [Husic and Pande \(2018\)](#) we approximate the continuous dynamics using closed-form calculations from the discrete Markov chain we have developed in the previous two steps. In particular, we estimate overall hitting probabilities using the standard “one-step analysis.” For any $k \in \{1, \dots, N_b\}$ and $i \in \{1, \dots, n-1\}$, let $u_k^{(i)}$ denote the probability of hitting $\partial A = \partial S_n$ before hitting $\partial \tilde{A} = \partial S_0$ if we start the discretized process at state k on ∂S_i . We can calculate our object of interest by solving the matrix difference equation

$$u^{(i)} = P^{(i,i+1)}u^{(i+1)} + P^{(i,i-1)}u^{(i-1)}, i = 1, \dots, n-1$$

with boundary conditions $u^{(0)} = \mathbf{0}, u^{(n)} = \mathbf{1}$, where $\mathbf{0}$ and $\mathbf{1}$ are vectors of all 0's and 1's. This gives the estimated hitting probability for each cluster. We then estimate the hitting probability of each point z_i by

$$h_{A, \tilde{A}^c}(z_i) = u_k^{(m)}, z_i \in \text{cluster } k \text{ on } \partial S \quad (2.4.2)$$

2.4.2.3 Concrete Choices for the Toy Model

To understand this algorithm more concretely, let us now use it to estimate $\text{cap}(A, \tilde{A})$ in the setup of the toy model. Here we can use the simple geometry and the exactly flat energy landscape of the toy model to our advantage. We follow the outline given

in Algorithm 2.4.2.

1. We pick $G = \dot{A}$ and $\tilde{G} = \tilde{A}$, where $\dot{A}$ and $\tilde{A}$ are as in our toy model (see Figure 2.1).
2. The surfaces areas can be easily calculated analytically as

$$|\partial G| = \frac{2\pi^{\frac{d}{2}}}{\Gamma(\frac{d}{2})} d_A^{d-1}, |\partial \tilde{G}| = \frac{2\pi^{\frac{d}{2}}}{\Gamma(\frac{d}{2})} \tilde{r}_A^{d-1}$$

3. The restrictions of the invariant measure $e^{-U(x)}dx$ on ∂G and $\partial \tilde{G}$ are simply uniform distributions on these two spheres, and we can generate exact samples on these two spheres. Since $\tilde{G} = \tilde{A}$ it is straightforward to use the probabilistic interpretation of h to see that $h_{A, \tilde{A}^c}(x) = 0, \forall x \in \partial \tilde{G}$, so we only need to apply Algorithm 2.4.3 to the case of $S = G = \dot{A}$. Following the notations used in Algorithm 2.4.3, assume we get N_p samples $z_1, \dots, z_{N_p}$, and the hitting probability vector $u^{(m)}$ for the N_b clusters on $\partial S_m = \partial G$, we have the estimates given in Equation 2.4.2
4. $h_{G, \tilde{G}^c}$ can be identified analytically as the harmonic function on $\tilde{G}/G$ Wendel (1980)

$$h_{G, \tilde{G}^c}(x) = \frac{1}{d_A^{2-d} - r_{\tilde{A}}^{2-d}} \|x - x_A\|^{2-d} - \frac{r_{\tilde{A}}^{2-d}}{d_A^{2-d} - r_{\tilde{A}}^{2-d}}$$

As a result, the gradient is given by

$$\nabla h_{G, \tilde{G}^c}(x) = \frac{2-d}{d_A^{2-d} - r_{\tilde{A}}^{2-d}} \|x - x_A\|^{1-d} \frac{x - x_A}{\|x - x_A\|}$$

Note that on both $\partial \tilde{G}$ and ∂G , the outward normal is given by $\mathbf{n}(x) = \frac{x - x_A}{\|x - x_A\|}$,

so

$$\mathbf{n}(x)^T \nabla h_{G, \tilde{G}^c}(x) = \frac{2-d}{d_A^{2-d} - r_{\tilde{A}}^{2-d}} \|x - x_A\|^{1-d} = \begin{cases} \frac{(2-d)r_{\tilde{A}}^{1-d}}{d_A^{2-d} - r_{\tilde{A}}^{2-d}} & \text{if } x \in \partial \tilde{G} \\ \frac{(2-d)d_A^{1-d}}{d_A^{2-d} - r_{\tilde{A}}^{2-d}} & \text{if } x \in \partial G \end{cases}$$

5. Finally, note that $U(x) = 0$ and $a(x) = I$ for $x \in \partial \tilde{G}, \partial G$. Plugging these into Equation (2.4.1), we obtain the approximation

$$\text{cap}(A, \tilde{A}) \approx \frac{2\pi^{\frac{d}{2}}}{\Gamma(\frac{d}{2})} \frac{d-2}{d_A^{2-d} - r_{\tilde{A}}^{2-d}} \frac{1}{N_p} \sum_{k=1}^{N_b} n_k u_k^{(m)} \quad (2.4.3)$$

where $n_k, k = 1, \dots, N_b$ is the number of samples in $z_1, \dots, z_{N_p}$ that belong to cluster k on $\partial S = \partial G = \partial \tilde{A}$

2.5 Numerical Experiments and Results

Recall that the central pieces of the CHop method are as follows:

- ε -flatness can be established for entropic barriers, and the hitting probability remains approximately a constant in the entropic barrier
- When ε -flatness holds, the hitting probability function can be well-approximated using CHop probabilities, which are defined in terms of the local property of capacities
- The capacities can be accurately and efficiently estimated by the CHop capacity estimation algorithm

In what follows, we present experimental results with our toy model, using both a flat and a nontrivial energy landscape, to verify each one of these aspects:

- We run multiple direct simulations from different starting points, to show that the hitting probabilities don't vary much, and verify the ε -flatness condition indeed holds for the parameters we use in the experiments.
- We estimate the CHop probabilities with capacities, given either analytically for the flat energy landscape, or numerically from our CHop capacity estimation algorithm for the nontrivial energy landscape, and compare with hitting probabilities from direct simulations, to show that the hitting probability function can be well-approximated using CHop probabilities.
- We compare the capacities estimated analytically and numerically for the flat energy landscape, to show the capacities can be accurately estimated using the CHop capacity estimation algorithm, and we compare the time it takes to estimate hitting probabilities using direct simulations with the time it takes to approximate hitting probabilities using CHop probabilities for the nontrivial energy landscape, to show the CHop method achieves high accuracy with considerably faster performance.

All the results presented in this section can be easily reproduced. For detailed instructions, please refer to https://github.com/StannisZhou/capacity_hopping.

2.5.1 Sanity Checks on Flat Energy Landscape

The first set of experiments are some basic sanity checks, with the goal of demonstrating the basic components of our algorithm. For these sanity checks, we work

with a flat energy landscape in $\mathbb{R}^5$, where quantities of interest can be obtained in closed form.

2.5.1.1 ε -flatness and CHop Probabilities

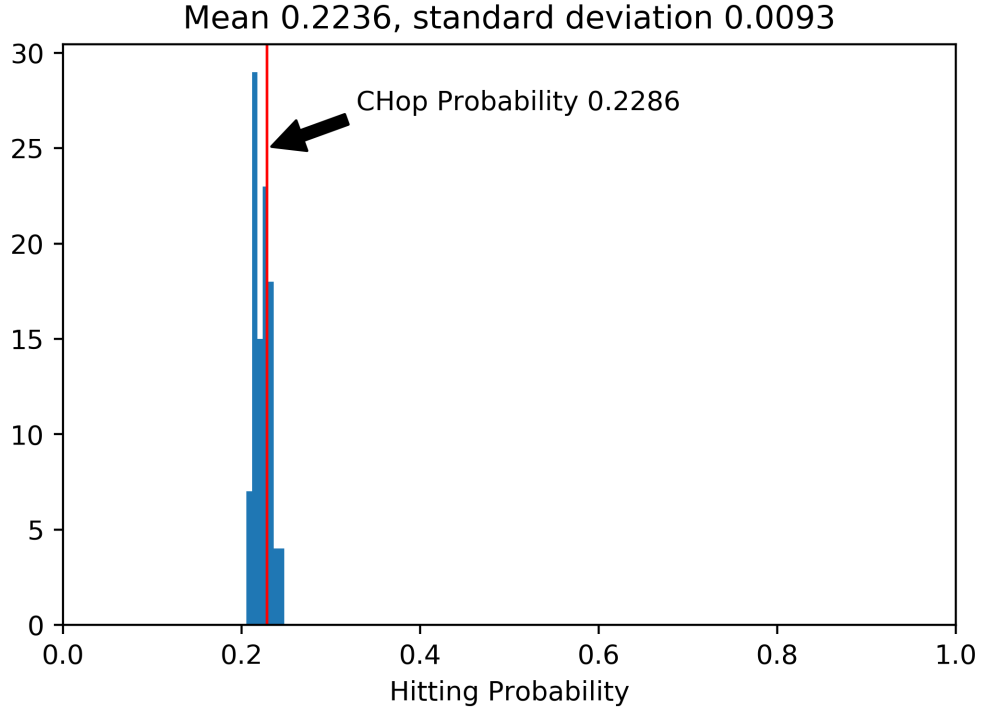
The first sanity check tests the CHop idea at its most basic level: that entropic barriers can be characterized by ε -flatness, and CHop probabilities calculated using capacities can be used to approximate global hitting probabilities. Here we use the toy model where $U(x) = 0, \forall x \notin A \cup B$. In other words, we have a flat energy landscape everywhere outside the targets. In this case, the capacity can be calculated analytically, and so the CHop probabilities can be computed in closed form:

$$p_A = \frac{\frac{1}{r_A^{2-d} - r_{\bar{A}}^{2-d}}}{\frac{1}{r_A^{2-d} - r_{\bar{A}}^{2-d}} + \frac{1}{r_B^{2-d} - r_{\bar{B}}^{2-d}}} \quad (2.5.1)$$

We test ε -flatness by looking at direct simulation results from different starting points and see if they give roughly the same hitting probabilities. We further test whether the CHop probabilities agree with direct simulation results. Theorem 2.3.3 shows that ε -flatness is valid, and CHop probabilities agree with global hitting probabilities in the limiting regime of small r_A, r_B , but it is not obvious whether the parameters we use lie in that regime. In particular, we look at the case $r_A = 0.05, r_{\bar{A}} = 0.1, r_B = 0.075, r_{\bar{B}} = 0.15$. In this case, we see that the ε -flatness condition indeed holds, and the CHop probability yields an estimate for the hitting probability that is consistent with direct simulations:

Hitting probabilities estimated with direct simulations ≈ 0.2236 . We run 2000 diffusion simulations at each one of the 100 randomly picked initial loca-

Figure 2.2: We picked 100 random locations for $X_0 \notin \mathcal{B}(0, 1) \setminus (\tilde{A} \cup \tilde{B})$. For each location we used 2000 simulations to estimate the probability of hitting A before B . The histogram of the 100 estimates is shown below. The ε -flatness condition holds with $\varepsilon = 0.0425$, and the CHop probability gives good approximation to this narrow range of hitting probabilities.



tions from $\mathcal{B}(0, 1) \setminus (\tilde{A} \cup \tilde{B})$. For each initial location we obtain an estimate of the hitting probability. The mean of these estimates is 0.2236, standard deviation 0.0093. We used time step of size $1e-05$. A histogram of the different hitting probabilities, as well as the CHop probability, is shown in Fig. 2.2. The hitting probabilities from different initial locations are in $[0.2055, 0.2480]$, and the ε -flatness condition holds with $\varepsilon = 0.0425$.

CHop probability 0.2286, which is a good approximation for the various hitting probabilities we obtained from direct simulations.

2.5.1.2 Capacity Estimation on Flat Energy Landscape

The second sanity check tests the capacity estimation algorithm employed by CHop. When the energy is flat, we can get an analytical formula for the capacity in the case of concentric spheres. For this test, we try to estimate the capacity $cap(A, \tilde{A})$, with $A = \mathcal{B}(0, r_A)$ and $\tilde{A} = \mathcal{B}(0, r_{\tilde{A}})$, and $U(x) = 0, \forall x \in \mathcal{B}(0, r_{\tilde{A}}) \setminus \mathcal{B}(0, r_A)$. In this case, the capacity is given analytically as

$$cap(A, \tilde{A}) = \frac{2\pi^{\frac{d}{2}}}{\Gamma(\frac{d}{2})} \frac{d-2}{r_A^{2-d} - r_{\tilde{A}}^{2-d}} \quad (2.5.2)$$

Take $r_A = 0.1, r_{\tilde{A}} = 0.4$, and $d_A = 0.2$. We can then compare the exact capacity for this case with the capacity given by our capacity estimation algorithm:

Exact value 0.080210. This is the capacity given by the analytical formula.

Estimate from capacity estimation algorithm 0.078846. To obtain this value we use the version of the capacity estimation algorithm described in the previous section, with parameter values given by

$$m = 2, n = 4, N_p = 100, N_b = 3, N_s = 1000$$

We used a 1e-06 time step.

2.5.2 Results on Nontrivial Energy Landscape

In this section, we present some results for the toy model with nontrivial energy landscape around the targets. In addition to testing ε -flatness and CHop probabilities, as we did in previous sanity checks, we further test the efficiency of our capacity estimation algorithm, to demonstrate that our algorithm can accurately estimate the global hitting probabilities while maintaining great advantage over the direct simulations in terms of speed.

We experimented with the toy model with $d = 5$ and

$$x_A = \begin{pmatrix} 0.5 \\ 0.6 \\ 0.0 \\ 0.0 \\ 0.0 \end{pmatrix}, r_A = 0.02, d_A = 0.05, r_{\bar{A}} = 0.1$$

for target A , and

$$x_B = \begin{pmatrix} -0.7 \\ 0.0 \\ 0.0 \\ 0.0 \\ 0.0 \end{pmatrix}, r_B = 0.04, d_B = 0.075, r_{\bar{B}} = 0.15$$

for target B . We refer the reader to appendix [E](#) for details on the energy function, as well as the actual parameters we used in our experiments.

2.5.2.1 ε -flatness and CHop Probabilities

We again start by testing the basic idea of CHop: we test our assumption of ε -flatness using direct simulations, to see the applicability of this condition when there is a nontrivial energy landscape. We further compare the CHop Probabilities with the hitting probabilities given by the direct simulations.

Hitting probabilities estimated with direct simulations: ≈ 0.8180 . We run 2000 diffusion simulations at each one of the 100 randomly picked initial locations from $\mathcal{B}(0, 1) \setminus (\tilde{A} \cup \tilde{B})$. For each initial location we obtain an estimate of the hitting probability. The mean of these estimates is 0.8180, standard deviation 0.0096. We used time step of size 1e-05. A histogram of the different hitting probabilities, as well as the CHop probability, is shown in Fig. 2.3. The hitting probabilities from different initial locations are in $[0.7980, 0.8450]$, and the ε -flatness condition holds with $\varepsilon = 0.0470$.

CHop probability 0.8274, which is in good agreement with the hitting probabilities we obtained from direct simulations. For the region around x_A , we used

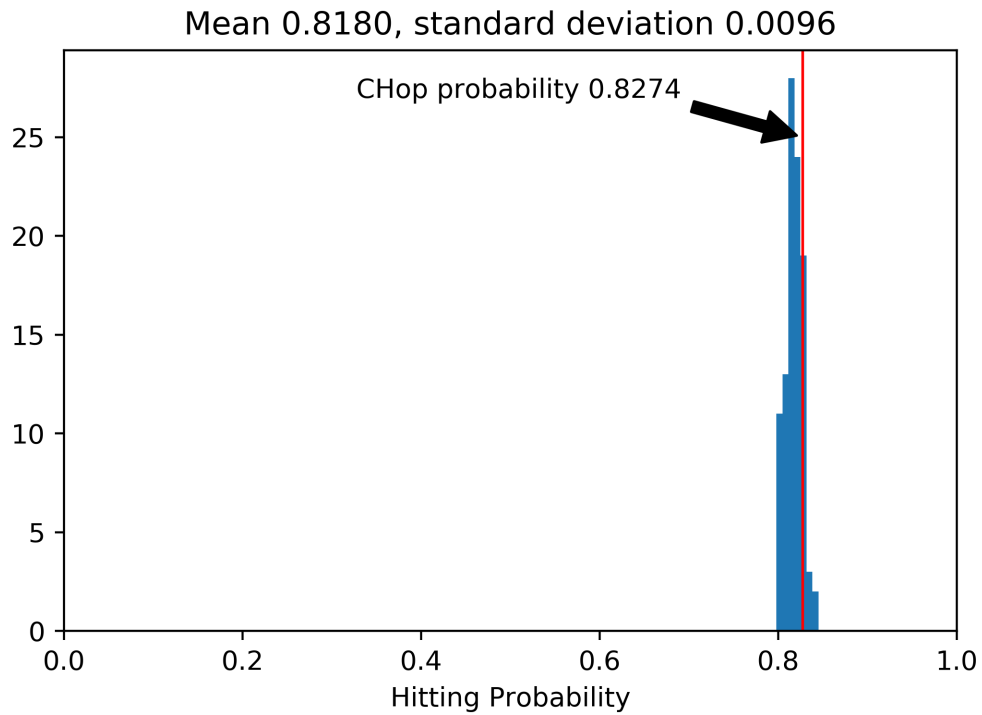
$$m = 2, n = 4, N_p = 5000, N_b = 10, N_s = 2000$$

For the region around x_B , we used

$$m = 2, n = 5, N_p = 5000, N_b = 10, N_s = 2000$$

For estimating both of these capacities, we used a 1e-07 time step.

Figure 2.3: We picked 100 random locations for $X_0 \notin \mathcal{B}(0, 1) \setminus (\tilde{A} \cup \tilde{B})$. For each location we used 2000 simulations to estimate the probability of hitting A before B . The histogram of the 100 estimates is shown below. The ε -flatness condition holds with $\varepsilon = 0.0470$, and the CHop probability gives good approximation to this narrow range of hitting probabilities.



2.5.2.2 Efficiency of the Capacity Estimation Algorithm

To test the efficiency of the algorithm, we compared the time it takes to estimate the hitting probabilities using direct simulations with the time our capacity estimation algorithm needs to estimate the capacities for all targets. Note that in order to make the direct simulations feasible for the 5-dimensional toy model we are working on, we made our best effort to make the direct simulation fast, including spherical acceleration in the flat region, JIT compilation to remove loop overhead, 24-CPU parallelization, and a relatively coarse time step.

Direct simulation 89177.87 seconds (3715.74 seconds on 24 CPUs). This quantity is an average over the amount of time it took to run 2000 simulations for each of the 100 different initial locations, using a relatively coarse time step of $1e-05$. We confirmed that there was no significant overhead due to the parallelization.

CHop speed 1055.55 seconds on a single CPU. Note that to get an accurate estimate, we need to employ a shorter time step of $1e-07$. But this doesn't constitute any computational burden because of the accelerations achieved by CHop.

This reflects an 85-fold acceleration. Furthermore note that the time-consuming elements of the capacity estimation algorithm are “embarrassingly parallelizable” and should be easy to further accelerate using parallelization. There was no need to do so in this case because we could easily run the CHop code on our laptops. However, larger scale problems would certainly benefit from this kind of acceleration.

The times reported above reflect a particular choice of parameters, but we found that the accuracy of the algorithm was fairly robust to choices of N_p , N_b and N_s .

For example, we performed another experiment with a $1\text{e-}06$ time step, parameters of $m = 2, n = 4, N_p = 3000, N_b = 5, N_s = 1000$ for estimating the A capacity, and parameters of $m = 2, n = 5, N_p = 3000, N_b = 5, N_s = 1000$ for estimating the B capacity. The result was an estimate of 0.8420. This still agrees well with the direct simulations (indeed, there were some initial conditions for which the direct simulation hitting probability estimates agreed with this quantity almost exactly). Using these parameters, total computation time was reduced to 119.35 seconds, reflecting a 750-fold acceleration.

2.6 Conclusions and Future Directions

Using a golf-course energy landscape with multiple targets as a prototypical example, this paper gives new theoretical insights into the commonly encountered entropic barriers in molecular dynamics simulations, from the perspective of hitting probabilities of the targets. A new concept called ε -flatness is proposed to capture the idea of the hitting probabilities being approximately a constant in a set, and ε -flatness is established for a characteristic entropic barrier, our toy model. Further connections are made between the global hitting probabilities of the entropic barriers and the capacities of local sets around the targets, to facilitate the easy understanding of entropic barriers. Inspired by these theoretical developments, we propose CHop as a general and practical method for overcoming entropic barriers, together with an efficient algorithm to deal with the central computational issue of CHop, namely the estimation of capacities.

Extensive numerical results demonstrate that CHop is highly effective in the presence of entropic barriers, both in terms of accuracy and speed, yet future work

remains. We see three directions as the most important in further developing the ideas set forth in this paper:

1. Move beyond the toy problem in this paper, and make CHop generally applicable to more realistic biomolecular systems. This involves identifying entropic barriers and establishing their ε -flatness in a more general setting, and thinking carefully about the geometry of more realistic biomolecular systems to estimate the capacities of the sets around the targets. Some preliminary discussions are given in this paper, but more work is needed in order to get a clear understanding.
2. Get results concerning hitting times, in addition to the hitting probabilities talked about in the paper, so that we can fully recover the kinetic information of the system.
3. Think more about ways to model the configuration space as a hierarchical, nested golf-course, with each target as a mini golf-course inside a larger one, and apply CHop recursively so as to develop a complete model for biomolecular systems, much like kinetic transition networks [Noé et al. \(2006\)](#), [Wales \(2006\)](#) and Markov State Models [Pande et al. \(2010\)](#), [Chodera and Noé \(2014\)](#), [Husic and Pande \(2018\)](#)

APPENDIX A

Stationary Reversible Diffusion Processes

In this section, we are going to make formal definitions of stationary reversible SDEs, and derive some properties that are central to the paper. Assume $\{X_t\}_{t \geq 0}$ is a stationary diffusion process in $\mathbb{R}^d$ with invariant measure μ . Define $\tau_S = \inf\{t \geq 0 : X_t \in S\}$, for $S \subset \mathbb{R}^d$. We first define the notion of reversibility for general diffusion processes. Our discussion here follows Section 4.6 of Pavliotis [Pavliotis \(2016\)](#).

Definition A.0.1. (Definition 4.3 of Pavliotis [Pavliotis \(2016\)](#)) A stationary stochastic process $\{X_t\}_{t \geq 0}$ is time-reversible if its law is invariant under time reversal: for every $T \in (0, +\infty)$, $\{X_t\}_{t \geq 0}$ and the time-reversed process $\{X_{T-t}\}_{t \geq 0}$ have the same distribution.

Recall that a diffusion process can be characterized by its generator. It turns out the reversibility of a stationary diffusion process can also be characterized by its generator. We have

Theorem A.0.2. (Theorem 4.5 of Pavliotis [Pavliotis \(2016\)](#)) A stationary diffusion process $\{X_t\}_{t \geq 0}$ in $\mathbb{R}^d$ with generator $\mathcal{L}$ and invariant measure μ is reversible if and only if its generator is self-adjoint in $L^2(\mathbb{R}^d; \mu)$.

For the rest of this section, we turn to the case where the stationary diffusion process $\{X_t\}_{t \geq 0}$ satisfies the SDE

$$dX_t = b(X_t)dt + \sigma(X_t)dW_t$$

where $b(\cdot) : \mathbb{R}^d \rightarrow \mathbb{R}^d$ and $\sigma(\cdot) : \mathbb{R}^d \rightarrow \mathbb{R}^{d \times m}$ are differentiable vector-valued and matrix-valued functions. Define $a(x) = \sigma(x)\sigma(x)^T$. We assume the invariant measure may be written as $\mu = e^{-U(x)}$, where $U(x)$ is a scalar function and can be thought of as a generalized potential.

Recall that, for this SDE, the generator $\mathcal{L}$ is given by

$$(\mathcal{L}f)(x) = b(x) \cdot \nabla f(x) + \frac{1}{2}a(x) : \nabla^2 f(x)$$

where $\cdot$ represents the inner product, $:$ represents the Frobenius inner product, and $\nabla^2 f(x)$ represents the Hessian matrix of the function $f(x)$, i.e.

$$(\mathcal{L}f)(x) = \sum_{i=1}^d b_i(x) \frac{\partial f(x)}{\partial x_i} + \frac{1}{2} \sum_{i=1}^d \sum_{j=1}^d a_{ij}(x) \frac{\partial^2 f(x)}{\partial x_i \partial x_j}$$

Using m to denote the Lebesgue measure, the adjoint of $\mathcal{L}$ in $L^2(\mathbb{R}^d; m)$ is the Fokker Planck operator

$$(\mathcal{L}^* f)(x) = -\nabla \cdot (b(x)f(x)) + \frac{1}{2} \nabla^2 : (a(x)f(x)) = -\sum_{i=1}^d \frac{\partial (b_i(x)f(x))}{\partial x_i} + \frac{1}{2} \sum_{i=1}^d \sum_{j=1}^d \frac{\partial^2 (a_{ij}(x)f(x))}{\partial x_i \partial x_j}$$

Note that we can further write the Fokker Planck operator $\mathcal{L}^*$ as $\mathcal{L}^* = \nabla \cdot J$, where

$$(Jf)(x) = -b(x)f(x) + \frac{1}{2} \nabla \cdot (a(x)f(x))$$

i.e.

$$(Jf)_i(x) = -b_i(x)f(x) + \frac{1}{2} \sum_{j=1}^d \frac{\partial (a_{ij}(x)f(x))}{\partial x_j}$$

In the case where $f(x)$ is a density function, $(Jf)(x)$ gives us the probability flux of the Markov process starting at the distribution $f(x)$.

Combining these basic formulas with Theorem [A.0.2](#), it's easy to see that the condition for the SDE being reversible is

$$(\mathcal{L}^* f e^{-U})(x) = e^{-U(x)} (\mathcal{L}f)(x), \forall f$$

Alternatively, we have the following proposition

Proposition A.0.3. (*Proposition 4.5 of Pavliotis [Pavliotis \(2016\)](#)*) X_t is reversible if and only if $(Je^{-U})(x)$ holds, or equivalently, $b(x) = \frac{1}{2}\nabla \cdot a(x) - \frac{1}{2}a(x)\nabla U(x)$.

We establish a useful property of stationary reversible SDEs.

Proposition A.0.4. *If X_t is reversible, then $(Jfe^{-U})(x) = \frac{1}{2}e^{-U(x)}a(x)\nabla f(x), \forall f$.*

Proof This result can be established by using the definition of J , the property $(Je^{-U})(x) = 0$, and straightforward calculations. We have

$$\begin{aligned}
 (Jfe^{-U})(x) &= -b(x)f(x)e^{-U(x)} + \frac{1}{2}\nabla \cdot (a(x)f(x)e^{-U(x)}) \\
 \nabla \cdot (a(x)f(x)e^{-U(x)}) &= f(x)\nabla \cdot (a(x)e^{-U(x)}) + e^{-U(x)}a(x)\nabla f(x) \\
 \implies (Jfe^{-U})(x) &= f(x) \left[-b(x)e^{-U(x)} + \frac{1}{2}\nabla \cdot (a(x)e^{-U(x)}) \right] + \frac{1}{2}e^{-U(x)}a(x)\nabla f(x) \\
 &= f(x)(Je^{-U})(x) + \frac{1}{2}e^{-U(x)}a(x)\nabla f(x) \\
 &= \frac{1}{2}e^{-U(x)}a(x)\nabla f(x)
 \end{aligned}$$

□

We next move on to the definition of Dirichlet form and capacity for stationary reversible SDEs. We have the following definitions

Definition A.0.5. Assume X_t is a stationary reversible SDE with invariant measure $\mu = e^{-U(x)}dx$, the Dirichlet form of X_t is given by

$$\mathcal{E}(f, g) = \frac{1}{2} \int_{\mathbb{R}^d} \nabla f(x)^T a(x) \nabla g(x) e^{-U(x)} dx$$

Note that in the case of an SDE in Ω with reflecting boundaries on $\partial\Omega$, the

stationary measure μ is restricted in Ω . In this case, the Dirichlet form is given by

$$\mathcal{E}(f, g) = \frac{1}{2} \int_{\Omega} \nabla f(x)^T a(x) \nabla g(x) e^{-U(x)} dx$$

So this definition also applies to SDEs with reflecting boundaries, which is the main object of interest in this paper. In what follows, w.l.o.g., we are going to use Ω to denote the state space of the process.

Definition A.0.6. Let $A \subset \tilde{A} \subset \Omega$ be non-empty. Assume we have an open set $D \subset \Omega$ for which $\partial D = \partial A \cup \partial \tilde{A}$ and is regular (in the sense of Theorem 7.17 in Bovier [Bovier and den Hollander \(2016\)](#)). Denote the Dirichlet form of the reversible stationary diffusion process X_t in Ω with invariant measure $\mu = e^{-U(x)} dx$ by $\mathcal{E}$. Let $\mathcal{H}_{A, \tilde{A}}$ be the space of functions f such that

1. $f \in H^1$, where H^1 is the Soblev space of weakly differentiable functions
2. $\mathcal{E}(f, f) < \infty$
3. $f \geq 1$ on ∂A and $f \leq 0$ on $\partial \tilde{A}$

The capacity $\text{cap}(A, \tilde{A})$ is defined as

$$\text{cap}(A, \tilde{A}) = \inf_{f \in \mathcal{H}_{A, \tilde{A}}} \mathcal{E}(f, f)$$

Before we state the basic properties of the capacity $\text{cap}(\tilde{A}, \tilde{A})$, we are also going to define the Dirichlet problem. This problem turns out to be just another perspective on the same minimization problem which appears in the definition of the capacity.

Definition A.0.7. Let $A, B \subset \Omega$ be non-empty and disjoint. Assume we have an open set $D \subset \Omega$ for which $\partial D = \partial A \cup \partial B$ and is regular. Given a stationary

diffusion process X_t in Ω with generator $\mathcal{L}$, the Dirichlet problem is defined as the partial differential equation

$$-(\mathcal{L}h)(x) = 0, \forall x \in D$$

with boundary conditions

$$h(x) = 1, x \in A$$

$$h(x) = 0, x \in B$$

Under certain regularity conditions (in the form of Theorem 7.17 in Bovier [Bovier and den Hollander \(2016\)](#)), the solution $h_{A,B}(x)$ exists and is unique, and has the probabilistic interpretation that

$$h_{A,B}(x) = \mathbb{P}(X_{\tau_{A \cup B}} \in A | X_0 = x)$$

i.e. $h_{A,B}(x)$ is the probability for X_t to hit A first before hitting B if we start the process at x . This is the celebrated *Dirichlet Principle*.

Armed with this definition, we are ready to state some basic properties of the capacity.

Proposition A.0.8. *The capacity $\text{cap}(A, \tilde{A})$ satisfies the following properties*

1. *(Probabilistic interpretation) If $\mathcal{H}_{A,\tilde{A}} \neq \emptyset$, then the infimum in the definition of $\text{cap}(A, \tilde{A})$ is attained uniquely at the solution to the corresponding Dirichlet problem $h_{A,\tilde{A}^c}$, i.e. $\text{cap}(A, \tilde{A}) = \mathcal{E}(h_{A,\tilde{A}^c}, h_{A,\tilde{A}^c})$.*
2. *(Additivity) The capacity is additive, i.e. if we have $A \subset \tilde{A} \subset \Omega$ and $B \subset \tilde{B} \subset \Omega$ with $\tilde{A} \cap \tilde{B} = \emptyset$, then*

Ω where $\tilde{A}$ and $\tilde{B}$ are disjoint, then we have

$$\text{cap}(A \cup B, \tilde{A} \cup \tilde{B}) = \text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})$$

Proof

1. Refer to Theorem 7.33 of Bovier [Bovier and den Hollander \(2016\)](#)
2. By definition, we have

$$\begin{aligned} h_{A, \tilde{A}^c}(x) &= 0, \forall x \in \tilde{A}^c \\ h_{B, \tilde{B}^c}(x) &= 0, \forall x \in \tilde{B}^c \end{aligned}$$

which implies

$$\begin{aligned} \nabla h_{A, \tilde{A}^c}(x) &= 0, \forall x \in \tilde{A}^c \\ \nabla h_{B, \tilde{B}^c}(x) &= 0, \forall x \in \tilde{B}^c \end{aligned}$$

Since $\tilde{A}$ and $\tilde{B}$ are disjoint, using the probabilistic interpretation, it's easy to see that

$$\begin{aligned} h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x) &= h_{A, \tilde{A}^c}(x), \forall x \in \tilde{A} \\ h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x) &= h_{B, \tilde{B}^c}(x), \forall x \in \tilde{B} \end{aligned}$$

Furthermore, we have $h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x) = 0, \forall x \in (\tilde{A} \cup \tilde{B})^c$, which implies

$$\nabla h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x) = 0, \forall x \in (\tilde{A} \cup \tilde{B})^c$$

Since X_t is reversible, using 1, we have

$$\begin{aligned}
& \text{cap}(A \cup B, \tilde{A} \cup \tilde{B}) \\
&= \mathcal{E}(h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}, h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}) \\
&= \frac{1}{2} \int_{\Omega} (\nabla h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x))^T a(x) \nabla h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x) e^{-U(x)} dx \\
&= \frac{1}{2} \int_{\tilde{A}} (\nabla h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x))^T a(x) \nabla h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x) e^{-U(x)} dx \\
&+ \frac{1}{2} \int_{\tilde{B}} (\nabla h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x))^T a(x) \nabla h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x) e^{-U(x)} dx \\
&+ \frac{1}{2} \int_{(\tilde{A} \cup \tilde{B})^c} (\nabla h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x))^T a(x) \nabla h_{A \cup B, (\tilde{A} \cup \tilde{B})^c}(x) e^{-U(x)} dx \\
&= \frac{1}{2} \int_{\tilde{A}} \left(\nabla h_{A, \tilde{A}^c}(x) \right)^T a(x) \nabla h_{A, \tilde{A}^c}(x) e^{-U(x)} dx \\
&+ \frac{1}{2} \int_{\tilde{B}} (\nabla h_{B, \tilde{B}^c}(x))^T a(x) \nabla h_{B, \tilde{B}^c}(x) e^{-U(x)} dx \\
&= \frac{1}{2} \int_{\Omega} \left(\nabla h_{A, \tilde{A}^c}(x) \right)^T a(x) \nabla h_{A, \tilde{A}^c}(x) e^{-U(x)} dx \\
&+ \frac{1}{2} \int_{\Omega} (\nabla h_{B, \tilde{B}^c}(x))^T a(x) \nabla h_{B, \tilde{B}^c}(x) e^{-U(x)} dx \\
&= \mathcal{E}(h_{A, \tilde{A}^c}(x), h_{A, \tilde{A}^c}(x)) + \mathcal{E}(h_{B, \tilde{B}^c}(x), h_{B, \tilde{B}^c}(x)) \\
&= \text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})
\end{aligned}$$

which proves the result.

□

APPENDIX B

Proof of the Main Theorem in Chapter 2

Proof of Theorem 2.3.3 Note that because of the probabilistic interpretation of the capacity, we have

$$\mathcal{E}(u, u) = \text{cap}(A, B^c) = \inf_{f \in \mathcal{H}_{A, B^c}} \mathcal{E}(f, f)$$

Recall that, since X_t is reversible, the Dirichlet form takes the form

$$\mathcal{E}(f, f) = \frac{1}{2} \int_{\Omega} (\nabla f(x))^T a(x) \nabla f(x) e^{-F(x)} dx$$

Intuitively, $\forall f \in \mathcal{H}_{A, B^c}$, the Dirichlet form $\mathcal{E}(f, f)$ measures the costs (in terms of suitable norm of the gradient of f) we need to pay to grow the function f from being 0 on ∂B to being 1 on ∂A . Because of the definition of u , $\forall f \in \mathcal{H}_{A, B^c}$, we have the upper bound

$$\mathcal{E}(u, u) \leq \mathcal{E}(f, f)$$

If we have a relatively tight upper-bound for $\mathcal{E}(u, u)$, it would place significant constraints on the function u itself, because of the ε -flatness of M . Since the function u can grow by at most ε within M , most of its growth from being 0 on ∂B to being 1 on ∂A would have to be done within $\tilde{A}/A$ and $\tilde{B}/B$. Increasing the amount of growth needed within a region by k times would approximately increase the contribution to the Dirichlet form within this region by k^2 times, so the value of u within M has to be carefully balanced. Being too far away from 0 would make

$$\frac{1}{2} \int_{\tilde{B}/B} (\nabla u(x))^T a(x) \nabla u(x) e^{-U(x)} dx$$

explode, while being too far away from 1 would make

$$\frac{1}{2} \int_{\tilde{A}/A} (\nabla u(x))^T a(x) \nabla u(x) e^{-U(x)} dx$$

explode. Both cases would violate the upper bound we have already established for $\mathcal{E}(u, u)$.

Formally, we first establish the upper bound by making use of $h_{A, \tilde{A}^c}$ and $h_{B, \tilde{B}^c}$. $\forall c \in (0, 1)$, define

$$\hat{u}_c(x) = \begin{cases} c, & \text{if } x \in M \\ (1-c)h_{A, \tilde{A}^c}(x) + c, & \text{if } x \in \tilde{A} \\ c(1-h_{B, \tilde{B}^c}(x)), & \text{if } x \in \tilde{B} \end{cases}$$

Obviously $\hat{u}_c \in \mathcal{H}_{A, B^c}$, which implies

$$\begin{aligned} & \mathcal{E}(u, u) \\ & \leq \mathcal{E}(\hat{u}_c, \hat{u}_c) \\ & = \frac{1}{2} \int_{\Omega} (\nabla \hat{u}_c(x))^T a(x) \nabla \hat{u}_c(x) e^{-U(x)} dx \\ & = \frac{1}{2} \int_{\tilde{A}} (\nabla \hat{u}_c(x))^T a(x) \nabla \hat{u}_c(x) e^{-U(x)} dx + \frac{1}{2} \int_{\tilde{B}} (\nabla \hat{u}_c(x))^T a(x) \nabla \hat{u}_c(x) e^{-U(x)} dx \\ & = \frac{1}{2} (1-c)^2 \int_{\tilde{A}} (\nabla h_{A, \tilde{A}^c}(x))^T a(x) \nabla h_{A, \tilde{A}^c}(x) e^{-U(x)} dx \\ & \quad + \frac{1}{2} c^2 \int_{\tilde{B}} (\nabla h_{B, \tilde{B}^c}(x))^T a(x) \nabla h_{B, \tilde{B}^c}(x) e^{-U(x)} dx \\ & = (1-c)^2 \text{cap}(A, \tilde{A}) + c^2 \text{cap}(B, \tilde{B}), \forall c \in (0, 1) \end{aligned}$$

Minimizing this w.r.t. c , we get an upper bound for $\mathcal{E}(u, u)$:

$$\mathcal{E}(u, u) \leq \frac{\text{cap}(A, \tilde{A}) \text{cap}(B, \tilde{B})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})}$$

Next, note that, because M is an ε -flat barrier for u , we have

$$\sup_{x \in M} u(x) - \inf_{x \in M} u(x) < \varepsilon$$

Define $m = \frac{1}{2}(\sup_{x \in M} u(x) + \inf_{x \in M} u(x))$, we have

$$|u(x) - m| < \frac{\varepsilon}{2}, \forall x \in M \quad (\text{B.0.1})$$

And we have $u(x) \geq m - \frac{\varepsilon}{2}, \forall x \in \partial \tilde{B}$, and $u(x) \leq m + \frac{\varepsilon}{2}, \forall x \in \partial \tilde{A}$. Since $u(x) = 1, \forall x \in A$, and $u(x) = 0, \forall x \in B$, define

$$\begin{aligned} u_A(x) &= \begin{cases} \frac{u(x) - m - \frac{\varepsilon}{2}}{1 - m - \frac{\varepsilon}{2}}, & x \in \tilde{A} \\ 0, & x \notin \tilde{A} \end{cases}, \text{ if } m + \frac{\varepsilon}{2} < 1 \\ u_B(x) &= \begin{cases} 1 - \frac{1}{m - \frac{\varepsilon}{2}} u(x), & x \in \tilde{B} \\ 0, & x \notin \tilde{B} \end{cases}, \text{ if } m > \frac{\varepsilon}{2} \end{aligned}$$

we have $u_A \in \mathcal{H}_{A, \tilde{A}}, u_B \in \mathcal{H}_{B, \tilde{B}}$, and

$$\begin{aligned} &\mathcal{E}(u, u) \\ &= \frac{1}{2} \int_{\Omega} (\nabla u(x))^T a(x) \nabla u(x) e^{-U(x)} dx \\ &\geq \frac{1}{2} \int_{\tilde{A}} (\nabla u(x))^T a(x) \nabla u(x) e^{-U(x)} dx + \frac{1}{2} \int_{\tilde{B}} (\nabla u(x))^T a(x) \nabla u(x) e^{-U(x)} dx \\ &\geq \left(1 - m - \frac{\varepsilon}{2}\right)^2 \mathcal{E}(u_A, u_A) \chi_{(0, 1 - \frac{\varepsilon}{2})}(m) + \left(m - \frac{\varepsilon}{2}\right)^2 \mathcal{E}(u_B, u_B) \chi_{(\frac{\varepsilon}{2}, 1)}(m) \\ &\geq \left(1 - m - \frac{\varepsilon}{2}\right)^2 \text{cap}(A, \tilde{A}) \chi_{(0, 1 - \frac{\varepsilon}{2})}(m) + \left(m - \frac{\varepsilon}{2}\right)^2 \text{cap}(B, \tilde{B}) \chi_{(\frac{\varepsilon}{2}, 1)}(m) \end{aligned}$$

Combining this with the upper bound we already establish, we have

$$\left(1 - m - \frac{\varepsilon}{2}\right)^2 \text{cap}(A, \tilde{A}) \chi_{(0, 1 - \frac{\varepsilon}{2})}(m) + \left(m - \frac{\varepsilon}{2}\right)^2 \text{cap}(B, \tilde{B}) \chi_{(\frac{\varepsilon}{2}, 1)}(m) \leq \frac{\text{cap}(A, \tilde{A}) \text{cap}(B, \tilde{B})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})}$$

In the typical case where $m \in (\frac{\varepsilon}{2}, 1 - \frac{\varepsilon}{2})$, the above constraint becomes the quadratic inequality

$$\left(1 - m - \frac{\varepsilon}{2}\right)^2 \text{cap}(A, \tilde{A}) + \left(m - \frac{\varepsilon}{2}\right)^2 \text{cap}(B, \tilde{B}) \leq \frac{\text{cap}(A, \tilde{A}) \text{cap}(B, \tilde{B})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})}$$

Solving this, we have

$$\begin{aligned}
m &\geq \frac{\text{cap}(A, \tilde{A}) \text{cap}(B, \tilde{B})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} + \frac{\varepsilon \text{cap}(B, \tilde{B}) - \text{cap}(A, \tilde{A})}{2 \text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} - \frac{2\sqrt{\text{cap}(A, \tilde{A}) \text{cap}(B, \tilde{B})}^{\frac{\varepsilon}{2}} \left(1 - \frac{\varepsilon}{2}\right)}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} \\
m &\leq \frac{\text{cap}(A, \tilde{A}) \text{cap}(B, \tilde{B})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} + \frac{\varepsilon \text{cap}(B, \tilde{B}) - \text{cap}(A, \tilde{A})}{2 \text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} + \frac{2\sqrt{\text{cap}(A, \tilde{A}) \text{cap}(B, \tilde{B})}^{\frac{\varepsilon}{2}} \left(1 - \frac{\varepsilon}{2}\right)}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})}
\end{aligned}$$

from which we deduce that $\left| m - \frac{\text{cap}(A, \tilde{A}) \text{cap}(B, \tilde{B})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} \right| \leq \frac{\varepsilon}{2} + \sqrt{\frac{\varepsilon}{2}}.$

In the extreme cases where m is outside $\left(\frac{\varepsilon}{2}, 1 - \frac{\varepsilon}{2}\right)$, in order for the inequality to hold, we would need to have $\frac{\text{cap}(A, \tilde{A})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})}$ to be very small when $m \leq \frac{\varepsilon}{2}$, and very big when $m \geq 1 - \frac{\varepsilon}{2}$, which would give us the desired bound. Because of symmetry, the proofs for these two cases are essentially the same. We take the case $m \leq \frac{\varepsilon}{2}$ as an example. In this case, we have the inequality

$$\left(1 - m - \frac{\varepsilon}{2}\right)^2 \text{cap}(A, \tilde{A}) \leq \frac{\text{cap}(A, \tilde{A}) \text{cap}(B, \tilde{B})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})}$$

which implies

$$\frac{\text{cap}(B, \tilde{B})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} \geq (1 - \varepsilon)^2$$

from which we have

$$\begin{aligned}
&\frac{\text{cap}(A, \tilde{A})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} \\
&= 1 - \frac{\text{cap}(B, \tilde{B})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} \\
&\leq 1 - (1 - \varepsilon)^2 \\
&= 2\varepsilon - \varepsilon^2
\end{aligned}$$

It's easy to see that when $\varepsilon \leq \frac{2}{9}$, we have

$$2\varepsilon - \varepsilon^2 - \left(\frac{\varepsilon}{2} + \sqrt{\frac{\varepsilon}{2}}\right) = \frac{3\varepsilon}{2} - \varepsilon^2 - \sqrt{\frac{\varepsilon}{2}} = \sqrt{\varepsilon} \left(\frac{3}{2}\sqrt{\varepsilon} - \varepsilon^{\frac{3}{2}} - \sqrt{\frac{1}{2}}\right) \leq \sqrt{\varepsilon} \left(\frac{3}{2}\frac{\sqrt{2}}{3} - \frac{\sqrt{2}}{2}\right) = 0$$

As a result, we have both $\frac{\text{cap}(A, \tilde{A})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} \leq \frac{\varepsilon}{2} + \sqrt{\frac{\varepsilon}{2}}$ and $m \leq \frac{\varepsilon}{2} \leq \frac{\varepsilon}{2} + \sqrt{\frac{\varepsilon}{2}}$, which implies

$$\left| m - \frac{\text{cap}(A, \tilde{A})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} \right| \leq \frac{\varepsilon}{2} + \sqrt{\frac{\varepsilon}{2}}$$

Putting this together with Equation [B.0.1](#), we have

$$\left| u(x) - \frac{\text{cap}(A, \tilde{A})}{\text{cap}(A, \tilde{A}) + \text{cap}(B, \tilde{B})} \right| \leq \varepsilon + \sqrt{\frac{\varepsilon}{2}}$$

□

APPENDIX C

Proof of the ε -flatness of the Toy Model in Chapter 2

Proof of Theorem 2.3.2 We use M_t to denote a simple Brownian motion trapped inside $\Omega = \mathcal{B}(0, 1)$ by reflecting boundaries. Define

$$\bar{u} = \mathbb{E}[u(Z)], \text{ where } Z \sim \text{Uniform}(\Omega)$$

$$\tilde{u}(x, t) = \mathbb{E}[u(M_t) | M_0 = x]$$

$$\tilde{\tau}_M = \inf\{t : M_t \in \mathcal{B}(x_A, d_A) \cup \mathcal{B}(x_B, d_B)\}$$

$$\tilde{\tau}_X = \inf\{t : X_t \in \mathcal{B}(x_A, d_A) \cup \mathcal{B}(x_B, d_B)\}$$

Note that, for our toy model X_t , since the energy landscape is flat outside $\mathcal{B}(x_A, d_A) \cup \mathcal{B}(x_B, d_B)$, $\tilde{\tau}_M | M_0 = x$ has the same distribution as $\tilde{\tau}_X | X_0 = x$, $\forall x \in \Omega / (\mathcal{B}(x_A, d_A) \cup \mathcal{B}(x_B, d_B))$. In particular, applying the Dynkin's formula, we have

$$u(x) = \mathbb{E}[u(X_{t \wedge \tilde{\tau}_X}) | X_0 = x] = \mathbb{E}[u(M_{t \wedge \tilde{\tau}_M}) | M_0 = x], \forall x \in \Omega / (\tilde{A} \cup \tilde{B})$$

Using the above, and the fact that $u(x) \in [0, 1], \forall x \in \Omega / (\tilde{A} \cup \tilde{B})$, it's easy to see

$$\begin{aligned} |u(x) - \tilde{u}(x, t)| &\leq \mathbb{E}[|u(M_{t \wedge \tilde{\tau}_M}) - u(M_t)| | M_0 = x] \\ &= \mathbb{E}[|u(M_{t \wedge \tilde{\tau}_M}) - u(M_t)| \chi_{(0, t)}(\tilde{\tau}_M) | M_0 = x] \\ &\leq \mathbb{P}(\tilde{\tau}_M < t | M_0 = x) \\ &= \mathbb{P}(\tilde{\tau}_X < t | X_0 = x) \end{aligned}$$

Using Lemma C.0.1, we also have

$$\begin{aligned}
|\tilde{u}(x, t) - \bar{u}| &\leq \left| \int_{\Omega} u(y) \mu_x^t(dy) - \int_{\Omega} u(y) \pi(dy) \right| \\
&\leq |\mu_x^t - \pi|(\Omega) \\
&= d(\mu_x^t, \pi) \\
&< \frac{2}{t}
\end{aligned}$$

where $|\mu_x^t - \pi|$ is the total variation of the signed measure $\mu_x^t - \pi$.

From Lemma C.0.2, for any fixed value of $d \geq 3$ and $r_{\tilde{A}}, r_{\tilde{B}}, \varepsilon > 0$, there exists a $c = c(d, r_{\tilde{A}}, r_{\tilde{B}}, \varepsilon)$ such that if $d_A, d_B < c$, then

$$\mathbb{P} \left(\tilde{\tau}_X < \frac{8}{\varepsilon} | X_0 \notin \tilde{A} \cup \tilde{B} \right) < \frac{\varepsilon}{4}$$

Likewise, for any fixed value of $d_A, r_{\tilde{A}}, d_B, r_{\tilde{B}}, \varepsilon > 0$, there exists a $c = c(d_A, r_{\tilde{A}}, d_B, r_{\tilde{B}}, \varepsilon)$, such that if the dimension $d \geq c$ then the same holds.

This implies that, $\forall x \notin \tilde{A} \cup \tilde{B}$

$$\begin{aligned}
|u(x) - \bar{u}| &\leq \left| u(x) - \tilde{u} \left(x, \frac{8}{\varepsilon} \right) \right| + \left| \tilde{u} \left(x, \frac{8}{\varepsilon} \right) - \bar{u} \right| \\
&\leq \mathbb{P} \left(\tilde{\tau}_X < \frac{8}{\varepsilon} | X_0 = x \right) + \frac{2}{\frac{8}{\varepsilon}} \\
&< \frac{\varepsilon}{4} + \frac{\varepsilon}{4} \\
&= \frac{\varepsilon}{2} \\
\implies \sup_{x, y \notin \tilde{A} \cup \tilde{B}} |u(x) - u(y)| &< \varepsilon
\end{aligned}$$

which is the desired result. $\square$

Lemma C.0.1. (*Uniform ergodicity*) Let $\Omega = \mathcal{B}(0, 1) \subset \mathbb{R}^d$, and use M_t to denote

a Brownian motion trapped by reflecting boundaries inside Ω . Use π to denote the uniform distribution within Ω , which is also the equilibrium distribution of M_t . Then the distribution $\mu_x^t = \mathbb{P}(M_t \in \cdot | M_0 = x)$ converges in total variation distance to π , and we can have the uniform bound

$$d(\mu_x^t, \pi) = \sup_{A \subset \Omega} |\mu_x^t(A) - \pi(A)| < \frac{2}{t}, \forall x \in \Omega, t > 0$$

Proof We adopt a coupling construction. Assume $M_0 = x$. Let Y_t denote another Brownian motion trapped by reflecting boundaries inside Ω , and assume $Y_0 \sim \pi$. Note that, due to the lack of any potential, we in fact have $Y_t \sim \pi, \forall t$. Define

$$\begin{aligned} \rho_1 &= \inf\{t : \|M_t\|_2 = \|Y_t\|_2\} \\ \alpha &= \frac{M_{\rho_1} - Y_{\rho_1}}{\|M_{\rho_1}\|_2 - \|Y_{\rho_1}\|_2} \end{aligned}$$

and further define

$$\tilde{Y}_t = \begin{cases} Y_t, t \leq \rho_1 \\ M_t - 2\alpha \langle \alpha, M_t \rangle, t \geq \rho_1 \end{cases}$$

Now take

$$\rho_2 = \inf\{t \geq \rho_1 : \langle \alpha, M_t \rangle = 0\}$$

and finally define

$$\tilde{\tilde{Y}}_t = \begin{cases} Y_t, t \leq \rho_1 \\ M_t - 2\alpha \langle \alpha, M_t \rangle, \rho_1 \leq t \leq \rho_2 \\ M_t, \rho_2 \leq t \end{cases}$$

Then $\tilde{\tilde{Y}}$ itself must carry the law of Brownian motion trapped inside of Ω with

$\tilde{\tilde{Y}}_0 \sim \pi$, and thus also satisfies $\tilde{\tilde{Y}}_t \sim \pi, \forall t > 0$. In particular, $\forall A \subset \Omega$,

$$\begin{aligned}
& |\mu_x^t(A) - \pi(A)| \\
&= |\mathbb{P}(M_t \in A) - \mathbb{P}(\tilde{\tilde{Y}}_t \in A)| \\
&= |\mathbb{E}[\chi_A(M_t) - \chi_A(\tilde{\tilde{Y}}_t)]| \\
&\leq \mathbb{P}(\rho_2 \geq t) \\
&\leq \frac{\mathbb{E}[\rho_2]}{t}
\end{aligned}$$

Thus to complete the proof of the lemma, it suffices to bound $\mathbb{E}[\rho_2]$. Use W_t to denote a d -dimensional Brownian motion. Define

$$\tau^W = \inf\{t : \|W_t\| \geq 1\}$$

Then

$$\mathbb{E}[\tau^W | W_0 = x] = \frac{1 - \|x\|_2^2}{d} \quad (\text{C.0.1})$$

This can be easily verified by writing down the corresponding Poisson equation for the expectation $f(x) = \mathbb{E}[\tau^W | W_0 = x]$.

$$\begin{aligned}
-\frac{1}{2}\Delta f(x) &= 1 \\
f(x) &= 0, \forall x \in \partial\mathcal{B}(0, 1)
\end{aligned}$$

With this result, we first bound $\mathbb{E}[\rho_1]$. Define $\rho_1^M = \inf\{t : \|M_t\|_2 = 1\}$ and $\rho_1^Y = \inf\{t : \|Y_t\|_2 = 1\}$. Since $\|X_t\|, \|Y_t\| \leq 1, \forall t$, it follows that $\rho_1 \leq \rho_1^M \wedge \rho_1^Y$. Applying Equation C.0.1, we have

$$\mathbb{E}[\rho_1] \leq \mathbb{E}[\rho_1^M] = \frac{1 - \|x\|_2^2}{d} \leq \frac{1}{d} < 1, \forall d \geq 3$$

To bound $\mathbb{E}[\rho_2 - \rho_1]$, we note that $\rho_2 - \rho_1$ is stochastically dominated by the time it takes a one-dimensional reflecting Brownian motion in $[-1, 1]$ to hit 0, starting at $\|x\|_2$. Using the reflection principle, we can see that this hitting time is equivalent to the first exit time from the interval $[-1, 1]$ of a one-dimensional Brownian motion starting at $1 - \|x\|_2$. Using Equation C.0.1, we have

$$\mathbb{E}[\rho_2 - \rho_1] \leq \frac{1 - (1 - \|x\|_2)^2}{1} \leq 1$$

Putting these together, we have

$$d(\mu_x^t, \pi) = \sup_{A \subset \Omega} |\mu_x^t(A) - \pi(A)| < \frac{2}{t}, \forall x \in \Omega, t > 0$$

□

Lemma C.0.2. *Define $\tilde{\tau}_X = \inf\{t : X_t \in \mathcal{B}(x_A, d_A) \cup \mathcal{B}(x_B, d_B)\}$. For any fixed value of $d \geq 3$ and $r_{\tilde{A}}, r_{\tilde{B}}, \varepsilon > 0$, there exists a $c = c(d, r_{\tilde{A}}, r_{\tilde{B}}, \varepsilon)$ such that if $d_A, d_B < c$, then*

$$\mathbb{P}\left(\tilde{\tau}_X < \frac{8}{\varepsilon} | X_0 \notin \tilde{A} \cup \tilde{B}\right) < \frac{\varepsilon}{4}$$

Likewise, for any fixed value of $d_A, r_{\tilde{A}}, d_B, r_{\tilde{B}}, \varepsilon > 0$, there exists a $c = c(d_A, r_{\tilde{A}}, d_B, r_{\tilde{B}}, \varepsilon)$, such that if the dimension $d \geq c$ then the same holds.

Proof First, observe that, for a given t ,

$$\begin{aligned} & \mathbb{P}(\tilde{\tau}_X < t | X_0 \notin \tilde{A} \cup \tilde{B}) \\ = & \mathbb{P}(X_{\tilde{\tau}_X} \in \mathcal{B}(x_A, d_A) | X_0 \notin \tilde{A} \cup \tilde{B}) \mathbb{P}(\tilde{\tau}_X < t | X_{\tilde{\tau}_X} \in \mathcal{B}(x_A, d_A), X_0 \notin \tilde{A} \cup \tilde{B}) \\ + & \mathbb{P}(X_{\tilde{\tau}_X} \in \mathcal{B}(x_B, d_B) | X_0 \notin \tilde{A} \cup \tilde{B}) \mathbb{P}(\tilde{\tau}_X < t | X_{\tilde{\tau}_X} \in \mathcal{B}(x_B, d_B), X_0 \notin \tilde{A} \cup \tilde{B}) \end{aligned}$$

If we can establish

$$\begin{aligned}\mathbb{P}\left(\tilde{\tau}_X < \frac{8}{\varepsilon} | X_{\tilde{\tau}_X} \in \mathcal{B}(x_A, d_A), X_0 \notin \tilde{A} \cup \tilde{B}\right) &< \frac{\varepsilon}{4} \\ \mathbb{P}\left(\tilde{\tau}_X < \frac{8}{\varepsilon} | X_{\tilde{\tau}_X} \in \mathcal{B}(x_B, d_B), X_0 \notin \tilde{A} \cup \tilde{B}\right) &< \frac{\varepsilon}{4}\end{aligned}$$

Then we have proved the original result. The proofs for these two inequalities are exactly the same. In what follows, we are going to prove

$$\mathbb{P}\left(\tilde{\tau}_X < \frac{8}{\varepsilon} | X_{\tilde{\tau}_X} \in \mathcal{B}(x_A, d_A), X_0 \notin \tilde{A} \cup \tilde{B}\right) < \frac{\varepsilon}{4}$$

Intuitively, this inequality is saying that, given we start at a point outside $\tilde{A}, \tilde{B}$ and hit $\mathcal{B}(x_A, d_A)$ before hitting $\mathcal{B}(x_B, d_B)$, the probability for the hitting time to be small is small. To prove this, we pick an additional $\tilde{d}_A \in (d_A, r_{\tilde{A}})$, and consider the trajectories of X_t which start at X_0 and hit $\mathcal{B}(x_A, d_A)$ before hitting $\mathcal{B}(x_B, d_B)$. In order for X_t to hit $\mathcal{B}(x_A, d_A)$, it has to hit $\mathcal{B}(x_A, \tilde{d}_A)$ first. Once it hits $\mathcal{B}(x_A, \tilde{d}_A)$, there are two possibilities: either it continues to hit $\mathcal{B}(x_A, d_A)$, or it comes out of $\mathcal{B}(x_A, r_{\tilde{A}})$. When d_A becomes small or the dimension d becomes large, the probability for the process to hit $\mathcal{B}(x_A, d_A)$ before coming out of $\mathcal{B}(x_A, r_{\tilde{A}})$ is going to zero. Therefore it would take the process a lot of attempts before finally hitting $\mathcal{B}(x_A, d_A)$. Each failed attempt is associated with a certain amount of hitting time. Because of the large number of failed attempts, the accumulated hitting time would be long with high probability.

Formally, we are going to look at the behaviour of X_t in $\mathcal{B}(x_A, r_{\tilde{A}}) \setminus \mathcal{B}(x_A, d_A)$ when we start on a point $x \in \partial\mathcal{B}(x_A, \tilde{d}_A)$. Define the hitting time

$$T = \inf\{t : X_t \in \mathcal{B}(x_A, d_A) \cup \mathcal{B}(x_A, r_{\tilde{A}})^c\}$$

We care about two conditional distributions: $T||X_T - x_A| = d_A$ and $T||X_T - x_A| = r_{\tilde{A}}$. Note that because of symmetry, the starting point x won't affect these two distributions, as long as it's on $\partial\mathcal{B}(x_A, \tilde{d}_A)$. It's easy to see that, if we have a sequence of random variables

$$T_1^{\text{out}}, \dots, T_n^{\text{out}}, \dots \sim T||X_T - x_A| = r_{\tilde{A}}$$

a random variable

$$T^{\text{in}} \sim T||X_T - x_A| = d_A$$

and a random variable

$$N \sim \text{Geometric}(p), p = \mathbb{P}(|X_T - x_A| = d_A | |X_0 - x_A| = \tilde{d}_A) = \frac{\tilde{d}_A^{2-d} - r_{\tilde{A}}^{2-d}}{d_A^{2-d} - r_{\tilde{A}}^{2-d}}$$

all independent of each other, $\forall t > 0$, we have

$$\mathbb{P}(\tilde{\tau}_X < t | X_{\tilde{\tau}_X} \in \mathcal{B}(x_A, d_A), X_0 \notin \tilde{A} \cup \tilde{B}) \leq \mathbb{P}\left(\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}} < t\right)$$

So we only need to establish bounds for $\mathbb{P}\left(\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}} < t\right)$. For this, we note that, $\forall s > 0$,

$$\begin{aligned} & \mathbb{E}\left[e^{-s[\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}}]}\right] \\ = & \mathbb{E}\left[e^{-sT^{\text{in}}} \prod_{n=0}^N e^{-sT_n^{\text{out}}}\right] \\ = & \mathbb{E}[e^{-sT^{\text{in}}}] \mathbb{E}[(\mathbb{E}[e^{-sT^{\text{out}}}])^N] \\ = & \mathbb{E}[e^{-sT^{\text{in}}}] \sum_{k=0}^{+\infty} p((1-p)\mathbb{E}[e^{-sT^{\text{out}}}])^k \\ = & \frac{p\mathbb{E}[e^{-sT^{\text{in}}}]}{1 - (1-p)\mathbb{E}[e^{-sT^{\text{out}}}]} \end{aligned}$$

where $T^{\text{out}} \sim T||X_T - x_A| = r_{\tilde{A}}$.

If we can bound $\mathbb{E} \left[e^{-s[\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}}]} \right]$, we would then be able to bound $\mathbb{P} \left(\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}} < t \right)$. To bound $\mathbb{E} \left[e^{-s[\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}}]} \right]$, we make use of the results in [Wendel \(1980\)](#). Taking $n = 0$ in equation (8) and (9) in [Wendel \(1980\)](#), we have

$$\begin{aligned} \mathbb{E}[e^{-sT^{\text{in}}}] &= \left(\frac{d_A}{\tilde{d}_A} \right)^h \frac{I_h(\nu r_{\tilde{A}})K_h(\nu \tilde{d}_A) - I_h(\nu \tilde{d}_A)K_h(\nu r_{\tilde{A}})}{I_h(\nu r_{\tilde{A}})K_h(\nu d_A) - I_h(\nu d_A)K_h(\nu r_{\tilde{A}})} \\ \mathbb{E}[e^{-sT^{\text{out}}}] &= \left(\frac{r_{\tilde{A}}}{\tilde{d}_A} \right)^h \frac{I_h(\nu d_A)K_h(\nu \tilde{d}_A) - I_h(\nu \tilde{d}_A)K_h(\nu d_A)}{I_h(\nu d_A)K_h(\nu r_{\tilde{A}}) - I_h(\nu r_{\tilde{A}})K_h(\nu d_A)} \end{aligned}$$

where $h = \frac{d-2}{2}$, $\nu = \sqrt{2s}$, and I_h, K_h are the modified Bessel functions of the first and second kind of order h . Using these formulas, it's easy to see that

$$\begin{aligned} &\mathbb{E} \left[e^{-s[\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}}]} \right] \\ &= \frac{p\mathbb{E}[e^{-sT^{\text{in}}}]}{1 - (1-p)\mathbb{E}[e^{-sT^{\text{out}}}]} \\ &= \frac{p \left(\frac{d_A}{\tilde{d}_A} \right)^h (I_h(\nu r_{\tilde{A}})K_h(\nu \tilde{d}_A) - I_h(\nu \tilde{d}_A)K_h(\nu r_{\tilde{A}}))}{(I_h(\nu r_{\tilde{A}})K_h(\nu d_A) - I_h(\nu d_A)K_h(\nu r_{\tilde{A}})) + (1-p) \left(\frac{r_{\tilde{A}}}{\tilde{d}_A} \right)^h (I_h(\nu d_A)K_h(\nu \tilde{d}_A) - I_h(\nu \tilde{d}_A)K_h(\nu d_A))} \end{aligned}$$

To establish the desired bounds, we need some elementary properties of the modified Bessel functions. We refer the reader to [DLMFnoa](#) for details. In particular, we have

$$\lim_{z \rightarrow 0} I_h(z) = 0 \text{ DLMFnoa equation 10.30.1}$$

$$\lim_{z \rightarrow 0} K_h(z) = \infty \text{ DLMFnoa equation 10.30.2}$$

$$I_h(z) = \left(\frac{1}{2}z \right)^h \sum_{k=0}^{\infty} \frac{\left(\frac{1}{4}z^2 \right)^k}{k! \Gamma(h+k+1)} \text{ DLMFnoa equation 10.25.2}$$

$$I_h(z) \sim \frac{1}{\sqrt{2\pi h}} \left(\frac{ez}{2h} \right)^h \text{ as } h \rightarrow \infty \text{ DLMFnoa equation 10.41.1}$$

$$K_h(z) \sim \sqrt{\frac{\pi}{2h}} \left(\frac{ez}{2h} \right)^{-h} \text{ as } h \rightarrow \infty \text{ DLMFnoa equation 10.41.2}$$

Using these properties, and the explicit formula for p , we can see that, when $d_A \rightarrow 0$, we have

$$p \rightarrow 0, \left(\frac{d_A}{\tilde{d}_A} \right)^h \rightarrow 0, I_h(\nu d_A) \rightarrow 0, K_h(\nu d_A) \rightarrow \infty$$

$I_h(\nu r_{\tilde{A}})K_h(\nu \tilde{d}_A) - I_h(\nu \tilde{d}_A)K_h(\nu r_{\tilde{A}})$ would remain a constant, so the numerator

$$p \left(\frac{d_A}{\tilde{d}_A} \right)^h (I_h(\nu r_{\tilde{A}})K_h(\nu \tilde{d}_A) - I_h(\nu \tilde{d}_A)K_h(\nu r_{\tilde{A}})) \rightarrow 0$$

For the denominator, we have

$$\begin{aligned} & (I_h(\nu r_{\tilde{A}})K_h(\nu d_A) - I_h(\nu d_A)K_h(\nu r_{\tilde{A}})) \\ & + (1-p) \left(\frac{r_{\tilde{A}}}{\tilde{d}_A} \right)^h (I_h(\nu d_A)K_h(\nu \tilde{d}_A) - I_h(\nu \tilde{d}_A)K_h(\nu d_A)) \\ & = \left(I_h(\nu r_{\tilde{A}}) - (1-p) \left(\frac{r_{\tilde{A}}}{\tilde{d}_A} \right)^h I_h(\nu \tilde{d}_A) \right) K_h(\nu d_A) \\ & + \left((1-p) \left(\frac{r_{\tilde{A}}}{\tilde{d}_A} \right)^h K_h(\nu \tilde{d}_A) - K_h(\nu r_{\tilde{A}}) \right) I_h(\nu d_A) \end{aligned}$$

Using DLMF [noa](#) equation 10.25.2, since we picked $\tilde{d}_A < r_{\tilde{A}}$, we have

$$\begin{aligned} & I_h(\nu r_{\tilde{A}}) \\ & = \left(\frac{1}{2} \nu r_{\tilde{A}} \right)^h \sum_{k=0}^{\infty} \frac{\left(\frac{1}{4} \nu^2 r_{\tilde{A}}^2 \right)^k}{k! \Gamma(h+k+1)} \\ & = \left(\frac{r_{\tilde{A}}}{\tilde{d}_A} \right)^h \left(\frac{1}{2} \nu \tilde{d}_A \right)^h \sum_{k=0}^{\infty} \left(\frac{r_{\tilde{A}}}{\tilde{d}_A} \right)^{2k} \frac{\left(\frac{1}{4} \nu^2 \tilde{d}_A^2 \right)^k}{k! \Gamma(h+k+1)} \\ & > \left(\frac{r_{\tilde{A}}}{\tilde{d}_A} \right)^h \left(\frac{1}{2} \nu \tilde{d}_A \right)^h \sum_{k=0}^{\infty} \frac{\left(\frac{1}{4} \nu^2 \tilde{d}_A^2 \right)^k}{k! \Gamma(h+k+1)} \\ & = \left(\frac{r_{\tilde{A}}}{\tilde{d}_A} \right)^h I_h(\nu \tilde{d}_A) \end{aligned}$$

Since $K_h(\nu d_A) \rightarrow \infty$ and $I_h(\nu d_A) \rightarrow 0$ as $d_A \rightarrow 0$, we have the denominator goes to

∞ , which implies

$$\lim_{d_A \rightarrow 0} \mathbb{E} \left[e^{-s[\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}}]} \right] = 0$$

When the dimension $d \rightarrow \infty$, we have $h = \frac{d-2}{2} \rightarrow \infty$ and

$$p \rightarrow 0, \left(\frac{d_A}{\tilde{d}_A} \right)^h \rightarrow 0$$

In this case, using DLMF^{noa} equations 10.41.1 and 10.41.2, we can see that $\mathbb{E} \left[e^{-s[\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}}]} \right]$ behaves like

$$\begin{aligned} & \frac{p \left(\frac{d_A}{\tilde{d}_A} \right)^h \frac{1}{2h} \left[\left(\frac{r_{\tilde{A}}}{d_A} \right)^h - \left(\frac{\tilde{d}_A}{r_{\tilde{A}}} \right)^h \right]}{\frac{1}{2h} \left[\left(\frac{r_{\tilde{A}}}{d_A} \right)^h - \left(\frac{d_A}{r_{\tilde{A}}} \right)^h \right] + (1-p) \left(\frac{r_{\tilde{A}}}{d_A} \right)^h \frac{1}{2h} \left[\left(\frac{d_A}{\tilde{d}_A} \right)^h - \left(\frac{\tilde{d}_A}{d_A} \right)^h \right]} \\ &= \frac{p \left(\frac{d_A}{\tilde{d}_A} \right)^h}{\frac{\left(\frac{r_{\tilde{A}}}{d_A} \right)^h - \left(\frac{d_A}{r_{\tilde{A}}} \right)^h}{\left(\frac{r_{\tilde{A}}}{d_A} \right)^h - \left(\frac{\tilde{d}_A}{r_{\tilde{A}}} \right)^h} + (1-p) \left(\frac{r_{\tilde{A}}}{d_A} \right)^h \left[\frac{\left(\frac{d_A}{\tilde{d}_A} \right)^h - \left(\frac{\tilde{d}_A}{d_A} \right)^h}{\left(\frac{r_{\tilde{A}}}{d_A} \right)^h - \left(\frac{\tilde{d}_A}{r_{\tilde{A}}} \right)^h} \right]} \end{aligned}$$

when h is large. Since $d_A < \tilde{d}_A < r_{\tilde{A}}$, we have

$$\begin{aligned} \lim_{h \rightarrow \infty} \frac{\left(\frac{r_{\tilde{A}}}{d_A} \right)^h - \left(\frac{d_A}{r_{\tilde{A}}} \right)^h}{\left(\frac{r_{\tilde{A}}}{d_A} \right)^h - \left(\frac{\tilde{d}_A}{r_{\tilde{A}}} \right)^h} &= \infty \\ \lim_{h \rightarrow \infty} \left(\frac{r_{\tilde{A}}}{\tilde{d}_A} \right)^h \left[\frac{\left(\frac{d_A}{\tilde{d}_A} \right)^h - \left(\frac{\tilde{d}_A}{d_A} \right)^h}{\left(\frac{r_{\tilde{A}}}{d_A} \right)^h - \left(\frac{\tilde{d}_A}{r_{\tilde{A}}} \right)^h} \right] &= 0 \end{aligned}$$

This implies that

$$\lim_{d \rightarrow \infty} \mathbb{E} \left[e^{-s[\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}}]} \right] = 0$$

The above arguments establish the desired bounds for $\mathbb{E} \left[e^{-s[\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}}]} \right]$ for both the case of $d_A \rightarrow 0$ and the case of $d \rightarrow \infty$. These bounds in turn give us bounds

for

$$\mathbb{P} \left(\sum_{n=0}^N T_n^{\text{out}} + T^{\text{in}} < t \right)$$

which leads to bounds for $\mathbb{P}(\tilde{\tau}_X < t | X_{\tilde{\tau}_X} \in \mathcal{B}(x_A, d_A), X_0 \notin \tilde{A} \cup \tilde{B})$. Using the same argument, we can also establish bounds for $\mathbb{P}(\tilde{\tau}_X < t | X_{\tilde{\tau}_X} \in \mathcal{B}(x_B, d_B), X_0 \notin \tilde{A} \cup \tilde{B})$. Combining these bounds, we can prove the desired result $\square$

APPENDIX D

Proof of the lemma in Chapter 2

Proof of Lemma 2.4.1 The key here is to use the divergence theorem

$$\int_{\Omega} \nabla \cdot F(x) dx = \int_{\partial\Omega} F(x)^T \mathbf{n}(x) dS$$

and the property of stationary reversible SDEs (Proposition A.0.4)

$$J(f(x)e^{-U(x)}) = \frac{1}{2}e^{-U(x)}a(x)\nabla f(x), \forall f$$

and the fact that $\mathcal{L}(h_{A,\tilde{A}^c}(x)) = 0$.

Using these, we have

$$\begin{aligned} \text{cap}(A, \tilde{A}) &= \int_{\tilde{A}/A} \nabla h_{A,\tilde{A}^c}(x)^T a(x) \nabla h_{A,\tilde{A}^c}(x) e^{-U(x)} dx \\ &= \int_{\partial\tilde{A}} h_{A,\tilde{A}^c}(x) e^{-U(x)} [a(x) \nabla h_{A,\tilde{A}^c}(x)]^T \mathbf{n}(x) dS \\ &\quad - \int_{\partial A} h_{A,\tilde{A}^c}(x) e^{-U(x)} [a(x) \nabla h_{A,\tilde{A}^c}(x)]^T \mathbf{n}(x) dS \\ &\quad - \int_{\tilde{A}/A} h_{A,\tilde{A}^c}(x) \nabla \cdot [a(x) \nabla h_{A,\tilde{A}^c}(x) e^{-U(x)}] dx \\ &= - \int_{\partial A} e^{-U(x)} [a(x) \nabla h_{A,\tilde{A}^c}(x)]^T \mathbf{n}(x) dS \\ &\quad - \int_{\tilde{A}/A} h_{A,\tilde{A}^c}(x) \nabla \cdot 2J(h_{A,\tilde{A}^c}(x) e^{-U(x)}) dx \\ &= - \int_{\partial A} e^{-U(x)} [a(x) \nabla h_{A,\tilde{A}^c}(x)]^T \mathbf{n}(x) dS \\ &\quad - \int_{\tilde{A}/A} h_{A,\tilde{A}^c}(x) 2\mathcal{L}^*(h_{A,\tilde{A}^c}(x) e^{-U(x)}) dx \\ &= - \int_{\partial A} e^{-U(x)} [a(x) \nabla h_{A,\tilde{A}^c}(x)]^T \mathbf{n}(x) dS \\ &\quad - 2 \int_{\tilde{A}/A} \mathcal{L}(h_{A,\tilde{A}^c}(x)) h_{A,\tilde{A}^c}(x) e^{-U(x)} dx \\ &= - \int_{\partial A} e^{-U(x)} [a(x) \nabla h_{A,\tilde{A}^c}(x)]^T \mathbf{n}(x) dS \end{aligned}$$

We further have

$$\begin{aligned}
0 &= \int_{G/A} \mathcal{L}(h_{A,\tilde{A}^c}(x)) e^{-U(x)} dx \\
&= \int_{G/A} \mathcal{L}^*(h_{A,\tilde{A}^c}(x)) e^{-U(x)} dx \\
&= \int_{G/A} \nabla \cdot \left(\frac{1}{2} e^{-U(x)} a(x) \nabla f(x) \right) dx \\
&= \frac{1}{2} \int_{\partial G} e^{-U(x)} [a(x) \nabla h_{A,\tilde{A}^c}(x)]^T \mathbf{n}(x) dS \\
&\quad - \frac{1}{2} \int_{\partial A} e^{-U(x)} [a(x) \nabla h_{A,\tilde{A}^c}(x)]^T \mathbf{n}(x) dS
\end{aligned}$$

This implies

$$\begin{aligned}
&\int_{\partial A} e^{-U(x)} [a(x) \nabla h_{A,\tilde{A}^c}(x)]^T \mathbf{n}(x) dS \\
&= \int_{\partial G} e^{-U(x)} [a(x) \nabla h_{A,\tilde{A}^c}(x)]^T \mathbf{n}(x) dS
\end{aligned}$$

which further implies

$$\begin{aligned}
& \text{cap}(A, \tilde{A}) \\
&= - \int_{\partial G} e^{-U(x)} [a(x) \nabla h_{A, \tilde{A}^c}(x)]^T \mathbf{n}(x) dS \\
&= \int_{\partial \tilde{G}} h_{G, \tilde{G}^c}(x) e^{-U(x)} [a(x) \nabla h_{A, \tilde{A}^c}(x)]^T \mathbf{n}(x) dS \\
&\quad - \int_{\partial G} h_{G, \tilde{G}^c}(x) e^{-U(x)} [a(x) \nabla h_{A, \tilde{A}^c}(x)]^T \mathbf{n}(x) dS \\
&= 2 \int_{\tilde{G}/G} h_{G, \tilde{G}^c}(x) \mathcal{L}^*(h_{A, \tilde{A}^c}(x) e^{-U(x)}) dx \\
&\quad + \int_{\tilde{G}/G} \nabla h_{G, \tilde{G}^c}(x)^T a(x) \nabla h_{A, \tilde{A}^c}(x) e^{-U(x)} dx \\
&= 2 \int_{\tilde{G}/G} h_{A, \tilde{A}^c}(x) \mathcal{L}^*(h_{G, \tilde{G}^c}(x) e^{-U(x)}) dx \\
&\quad + \int_{\tilde{G}/G} \nabla h_{G, \tilde{G}^c}(x)^T a(x) \nabla h_{A, \tilde{A}^c}(x) e^{-U(x)} dx \\
&= \int_{\partial \tilde{G}} h_{A, \tilde{A}^c}(x) e^{-U(x)} [a(x) \nabla h_{G, \tilde{G}^c}(x)]^T \mathbf{n}(x) dS \\
&\quad - \int_{\partial G} h_{A, \tilde{A}^c}(x) e^{-U(x)} [a(x) \nabla h_{G, \tilde{G}^c}(x)]^T \mathbf{n}(x) dS
\end{aligned}$$

which proves the desired result. $\square$

APPENDIX E

Details on Energy Function used in Chapter 2

For the energy function, we hand-designed two different kinds of landscape: random well energy, which we use for the region around target A , and random crater energy, which we use for the region around target B . The basic components of these energy functions are the well component, given by

$$F_w(x|d_w, r) = -\frac{d_w}{r^4}(\|x - x_A\|_2^4 - 2r^2\|x - x_A\|_2^2) - d_w \quad (\text{E.0.1})$$

where d_w gives the depth of the well; the crater component, given by

$$F_c(x|d_c, h, r) = \frac{d_c}{3b^2r^4 - r^6}(2\|x - x_B\|_2^2 - 3(b^2 + r^2)\|x - x_B\|_2^4 + 6b^2r^2\|x - x_B\|_2^2) - d_c \quad (\text{E.0.2})$$

where d_c and h give the depth and the height of the crater, respectively, and

$$b^2 = -\frac{1}{3d}(-3d_cr^2 + C + \frac{\Delta_0}{C}) \quad (\text{E.0.3})$$

with

$$C = 3r^2\sqrt[3]{d_ch(d_c + \sqrt{d_c(d_c + h)})}$$

$$\Delta_0 = -9d_chr^4$$

and finally a random component, given by

$$F_r(x|\mu, \sigma) = \sum_{i=1}^m \prod_{j=1}^d \exp\left(-\frac{(x_j - \mu_{ij})^2}{2\sigma_{ij}^2}\right) \quad (\text{E.0.4})$$

where $\mu = (\mu_{ij})_{m \times d}$ and $\sigma = (\sigma_{ij})_{m \times d}$, with $\mu_i = (\mu_{i1}, \dots, \mu_{id})$, $i = 1, \dots, m$ being the locations of m Gaussian random bumps in the region around the targets, and σ_{ij} , $i = 1, \dots, m$, $j = 1, \dots, d$ gives the corresponding standard deviations.

To make sure the energy function is continuous, and the different components of

the energy function are balanced, we introduce a mollifier, given by

$$F_m(x|x_0, r) = \exp\left(-\frac{r}{r - \|x - x_0\|_2^{20}}\right) \quad (\text{E.0.5})$$

where $x_0 = x_A, r = d_A$ or $x_0 = x_B, r = d_B$, depending on which target we are working with, and a rescaling of the random component, which is given by $0.1 * d_w$ if we are perturbing the well component, and $0.1 * (d_c + h)$ if we are perturbing the crater component.

Intuitively, for the well component, we use a 4th order polynomial to get a well-like energy landscape around the target that is continuous and differentiable at the boundary. Similarly, for the crater component, we use a 6th order polynomial to get a crater-like energy landscape around the target that is also continuous and differentiable at the boundary. For the random component, we are essentially placing a number of Gaussian bumps around the target. And for the mollifier, we are designing the function such that it's almost exactly 1 around the target, until it comes to the outer boundary, when it transitions smoothly and swiftly to 0. To summarize, given parameters d_w, d_c, h and random bumps μ_A, μ_B with $\mu_i^A \in \dot{A} \setminus A, i = 1, \dots, m_A, \mu_i^B \in \dot{B} \setminus B, i = 1, \dots, m_B$, and the corresponding standard deviations σ^A, σ^B with $\sigma_{ij}^A, i = 1, \dots, m_A, j = 1, \dots, d, \sigma_{ij}^B, i = 1, \dots, m_B, j = 1, \dots, d$, the energy function we used in the experiments is given by

$$F(x) = F_w(x|d_w, d_A) + 0.1 \times d_w \times F_m(x|x_A, d_A) + F_r(x|\mu^A, \sigma^A), \forall x \in \dot{A} \setminus A \quad (\text{E.0.6})$$

$$F(x) = F_c(x|d_c, h, d_B) + 0.1 \times (d_c + h) \times F_m(x|x_B, d_B) + F_r(x|\mu^B, \sigma^B), \forall x \in \dot{B} \setminus B \quad (\text{E.0.7})$$

In our actual experiments, we used

$$d_w = 10.0, d_c = 6.0, h = 1.0, \sigma_{ij}^A = \sigma_{k,l}^B = 0.01, \forall i, j, k, l$$

Bibliography

- NIST Digital Library of Mathematical Functions. <http://dlmf.nist.gov/>. URL <http://dlmf.nist.gov/>.
- S. F. Altschul and B. W. Erickson. Significance of nucleotide sequence alignments: a method for random sequence permutation that preserves dinucleotide and codon usage. *Mol. Biol. Evol.*, 2(6):526–538, Nov. 1985. URL <https://www.ncbi.nlm.nih.gov/pubmed/3870875>.
- M. Andrec, A. K. Felts, E. Gallicchio, and R. M. Levy. Protein folding pathways from replica exchange simulations and a kinetic network model. *Proc. Natl. Acad. Sci. U. S. A.*, 102(19):6801–6806, May 2005. URL <http://www.pnas.org/content/102/19/6801>.
- D. Aristoff, J. Bello-Rivas, and R. Elber. A Mathematical Framework for Exact Milestoning. *Multiscale Model. Simul.*, 14(1):301–322, Jan. 2016. URL <https://doi.org/10.1137/15M102157X>.
- D. Baker and D. A. Agard. Kinetics versus thermodynamics in protein folding. *Biochemistry*, 33(24):7505–7509, June 1994. URL <https://www.ncbi.nlm.nih.gov/pubmed/8011615>.
- E. B. Baum. Intractable Computations without Local Minima. *Phys. Rev. Lett.*, 57(21):2764–2767, Nov. 1986. URL <https://link.aps.org/doi/10.1103/PhysRevLett.57.2764>.
- J. M. Bello-Rivas and R. Elber. Exact milestoning. *J. Chem. Phys.*, 142(9):094102, Mar. 2015. URL <https://doi.org/10.1063/1.4913399>.
- P. G. Bolhuis, D. Chandler, C. Dellago, and P. L. Geissler. Transition path sampling: throwing ropes over rough mountain passes, in the dark. *Annu. Rev. Phys. Chem.*, 53:291–318, 2002. URL <http://dx.doi.org/10.1146/annurev.physchem.53.082301.113146>.
- A. Bovier and F. den Hollander. *Metastability: A Potential-Theoretic Approach*. Springer, Feb. 2016. URL <https://market.android.com/details?id=book-KoyRCwAAQBAJ>.
- D. Chandler. Statistical mechanics of isomerization dynamics in liquids and the transition state approximation. *J. Chem. Phys.*, 68(6):2959–2970, Mar. 1978. URL <https://aip.scitation.org/doi/abs/10.1063/1.436049>.

- J. D. Chodera and F. Noé. Markov state models of biomolecular conformational dynamics. *Curr. Opin. Struct. Biol.*, 25:135–144, Apr. 2014. URL <http://dx.doi.org/10.1016/j.sbi.2014.04.002>.
- M. Christen and W. F. van Gunsteren. On searching in, sampling of, and dynamically moving through conformational space of biomolecular systems: A review. *J. Comput. Chem.*, 29(2):157–166, Jan. 2008. URL <http://doi.wiley.com/10.1002/jcc.20725>.
- C. Dellago, P. G. Bolhuis, F. S. Csajka, and D. Chandler. Transition path sampling and the calculation of rate constants. *J. Chem. Phys.*, 108(5):1964–1977, Feb. 1998. URL <https://doi.org/10.1063/1.475562>.
- R. O. Dror, R. M. Dirks, J. P. Grossman, H. Xu, and D. E. Shaw. Biomolecular Simulation: A Computational Microscope for Molecular Biology. *Annu. Rev. Biophys.*, 41(1):429–452, May 2012. URL <https://doi.org/10.1146/annurev-biophys-042910-155245>.
- W. E. and E. Vanden-Eijnden. Towards a Theory of Transition Paths. *J. Stat. Phys.*, 123(3):503, May 2006. URL <https://doi.org/10.1007/s10955-005-9003-9>.
- W. E. and E. Vanden-Eijnden. Transition-path theory and path-finding algorithms for the study of rare events. *Annu. Rev. Phys. Chem.*, 61:391–420, 2010. URL <http://dx.doi.org/10.1146/annurev.physchem.040808.090412>.
- H. Eyring. The Activated Complex in Chemical Reactions. *J. Chem. Phys.*, 3(2):107–115, Feb. 1935. URL <https://doi.org/10.1063/1.1749604>.
- W. M. Fitch. Random sequences. *J. Mol. Biol.*, 163(2):171–176, Jan. 1983. URL <https://www.ncbi.nlm.nih.gov/pubmed/6842586>.
- C. Flamm and I. L. Hofacker. Beyond energy minimization: approaches to the kinetic folding of RNA. *Monatsh. Chem.*, 139(4):447–457, Apr. 2008. URL <https://link.springer.com/article/10.1007/s00706-008-0895-3>.
- D. E. Goldberg. *Genetic Algorithms in Search, Optimization and Machine Learning*. Addison-Wesley Longman Publishing Co., Inc., Boston, MA, USA, 1st edition, 1989. URL <https://dl.acm.org/citation.cfm?id=534133>.
- J. M. Goldberg and R. L. Baldwin. A specific transition state for S-peptide combining with folded S-protein and then refolding. *Proc. Natl. Acad. Sci. U. S. A.*, 96(5):2019–2024, Mar. 1999. URL <https://www.ncbi.nlm.nih.gov/pubmed/10051587>.
- R. R. Gutell, A. Power, G. Z. Hertz, E. J. Putz, and G. D. Stormo. Identifying constraints on the higher-order structure of RNA: continued development and application of comparative sequence analysis methods. *Nucleic Acids Res.*, 20(21):5785–5795, Nov. 1992. URL <https://www.ncbi.nlm.nih.gov/pubmed/1454539>.
- P. G. Higgs. RNA secondary structure: physical and computational aspects. *Q. Rev. Biophys.*, 33(3):199–253, Aug. 2000. URL <https://www.ncbi.nlm.nih.gov/pubmed/11191843>.

- A. Hospital, J. R. Goñi, M. Orozco, and J. L. Gelpí. Molecular dynamics simulations: advances and applications. *Adv. Appl. Bioinform. Chem.*, 8:37–47, Nov. 2015. URL <http://dx.doi.org/10.2147/AABC.S70333>.
- X. Huang, G. R. Bowman, S. Bacallado, and V. S. Pande. Rapid equilibrium sampling initiated from nonequilibrium data. *Proc. Natl. Acad. Sci. U. S. A.*, 106(47):19765–19769, Nov. 2009. URL <http://www.pnas.org/content/106/47/19765>.
- X. Huang, Y. Yao, G. R. Bowman, J. Sun, L. J. Guibas, G. Carlsson, and V. S. Pande. Constructing multi-resolution Markov State Models (MSMs) to elucidate RNA hairpin folding mechanisms. *Pac. Symp. Biocomput.*, pages 228–239, 2010. URL <https://www.ncbi.nlm.nih.gov/pubmed/19908375>.
- B. E. Husic and V. S. Pande. Markov State Models: From an Art to a Science. *J. Am. Chem. Soc.*, 140(7):2386–2396, Feb. 2018. URL <https://doi.org/10.1021/jacs.7b12191>.
- Jackson Loper*, **Guangyao Zhou***, and S. Geman. Overcoming Entropic Barriers by Capacity Hopping. *Manuscript available upon request*, 2018.
- M. Jacob, M. Geeves, G. Holtermann, and F. X. Schmid. Diffusional barrier crossing in a two-state protein folding reaction. *Nat. Struct. Biol.*, 6(10):923–926, Oct. 1999. URL <http://dx.doi.org/10.1038/13289>.
- M. Jiang, J. Anderson, J. Gillespie, and M. Mayne. uShuffle: a useful tool for shuffling biological sequences while preserving the k-let counts. *BMC Bioinformatics*, 9:192, Apr. 2008. URL <http://dx.doi.org/10.1186/1471-2105-9-192>.
- D. Kandel, Y. Matias, R. Unger, and P. Winkler. Shuffling biological sequences. *Discrete Appl. Math.*, 71(1):171–185, Dec. 1996. URL <http://www.sciencedirect.com/science/article/pii/S0166218X97814564>.
- S. Kirkpatrick, C. D. Gelatt, and M. P. Vecchi. Optimization by Simulated Annealing. *Science*, 220(4598):671–680, May 1983. URL <http://science.sciencemag.org/content/220/4598/671>.
- K. Klenin, B. Strodel, D. J. Wales, and W. Wenzel. Modelling proteins: conformational sampling and reconstruction of folding kinetics. *Biochim. Biophys. Acta*, 1814(8):977–1000, Aug. 2011. URL <http://dx.doi.org/10.1016/j.bbapap.2010.09.006>.
- J. T. Y. Kung, D. Colognori, and J. T. Lee. Long noncoding RNAs: past, present, and future. *Genetics*, 193(3):651–669, Mar. 2013. URL <http://dx.doi.org/10.1534/genetics.112.146704>.
- C. Levinthal. How to fold graciously. *Mossbauer spectroscopy in biological systems*, 67:22–24, 1969.
- T. Z. Lwin and R. Luo. Overcoming entropic barrier with coupled sampling at dual resolutions. *J. Chem. Phys.*, 123(19):194904, Nov. 2005. URL <http://dx.doi.org/10.1063/1.2102871>.

- J. Machta. Strengths and weaknesses of parallel tempering. *Phys. Rev. E*, 80(5):056706, Nov. 2009. URL <https://link.aps.org/doi/10.1103/PhysRevE.80.056706>.
- D. H. Mathews, J. Sabina, M. Zuker, and D. H. Turner. Expanded sequence dependence of thermodynamic parameters improves prediction of RNA secondary structure. *J. Mol. Biol.*, 288(5):911–940, May 1999. URL <http://dx.doi.org/10.1006/jmbi.1999.2700>.
- J. A. McCammon, B. R. Gelin, and M. Karplus. Dynamics of folded proteins. *Nature*, 267:585, June 1977. URL <http://dx.doi.org/10.1038/267585a0>.
- T. C. B. McLeish. Protein Folding in High-Dimensional Spaces: Hypergutters and the Role of Nonnative Interactions. *Biophys. J.*, 88(1):172–183, Jan. 2005. URL <http://www.sciencedirect.com/science/article/pii/S0006349505730964>.
- S. Mohan, C. Hsiao, H. VanDeusen, R. Gallagher, E. Krohn, B. Kalahar, R. M. Wartell, and L. D. Williams. Mechanism of RNA Double Helix-Propagation at Atomic Resolution. *J. Phys. Chem. B*, 113(9):2614–2623, Mar. 2009. URL <http://dx.doi.org/10.1021/jp8039884>.
- D. Mondal and D. S. Ray. Diffusion over an entropic barrier: non-Arrhenius behavior. *Phys. Rev. E Stat. Nonlin. Soft Matter Phys.*, 82(3 Pt 1):032103, Sept. 2010. URL <http://dx.doi.org/10.1103/PhysRevE.82.032103>.
- K. V. Morris and J. S. Mattick. The rise of regulatory RNA. *Nat. Rev. Genet.*, 15(6):423–437, June 2014. URL <http://dx.doi.org/10.1038/nrg3722>.
- A. N. Naganathan and V. Muñoz. Scaling of Folding Times with Protein Size. *J. Am. Chem. Soc.*, 127(2):480–481, Jan. 2005. URL <https://doi.org/10.1021/ja044449u>.
- F. Noé, D. Krachtus, J. C. Smith, and S. Fischer. Transition Networks for the Comprehensive Characterization of Complex Conformational Change in Proteins. *J. Chem. Theory Comput.*, 2(3):840–857, May 2006. URL <https://doi.org/10.1021/ct050162r>.
- V. S. Pande, K. Beauchamp, and G. R. Bowman. Everything you wanted to know about Markov State Models but were afraid to ask. *Methods*, 52(1):99–105, Sept. 2010. URL <http://dx.doi.org/10.1016/j.ymeth.2010.06.002>.
- G. A. Pavliotis. Stochastic Processes and Applications — SpringerLink. <http://link.springer.com/content/pdf/10.1007/978-1-4939-1323-7.pdf>, 2016. URL <http://link.springer.com/content/pdf/10.1007/978-1-4939-1323-7.pdf>. Accessed: 2017-11-17.
- D. Perez, B. P. Uberuaga, Y. Shim, J. G. Amar, and A. F. Voter. Chapter 4 Accelerated Molecular Dynamics Methods: Introduction and Recent Developments. In R. A. Wheeler, editor, *Annual Reports in Computational Chemistry*, volume 5, pages 79–98. Elsevier, Jan. 2009. URL <http://www.sciencedirect.com/science/article/pii/S1574140009005040>.

- K. W. Plaxco and D. Baker. Limited internal friction in the rate-limiting step of a two-state protein folding reaction. *Proc. Natl. Acad. Sci. U. S. A.*, 95(23):13591–13596, Nov. 1998. URL <http://www.pnas.org/content/95/23/13591>.
- D. Pörschke. Model calculations on the kinetics of oligonucleotide double helix coil transitions. Evidence for a fast chain sliding reaction. *Biophys. Chem.*, 2(2):83–96, Aug. 1974a. URL <https://www.ncbi.nlm.nih.gov/pubmed/4433687>.
- D. Pörschke. A direct measurement of the unzipping rate of a nucleic acid double helix. *Biophys. Chem.*, 2(2):97–101, Aug. 1974b. URL <https://www.ncbi.nlm.nih.gov/pubmed/4433688>.
- D. Pörschke. Elementary steps of base recognition and helix-coil transitions in nucleic acids. *Mol. Biol. Biochem. Biophys.*, 24:191–218, 1977. URL <https://www.ncbi.nlm.nih.gov/pubmed/904614>.
- D. Reguera and J. M. Rubí. Kinetic equations for diffusion in the presence of entropic barriers. *Phys. Rev. E Stat. Nonlin. Soft Matter Phys.*, 64(6 Pt 1):061106, Dec. 2001. URL <http://dx.doi.org/10.1103/PhysRevE.64.061106>.
- H. A. Scheraga, M. Khalili, and A. Liwo. Protein-folding dynamics: overview of molecular simulation techniques. *Annu. Rev. Phys. Chem.*, 58:57–83, 2007. URL <http://dx.doi.org/10.1146/annurev.physchem.58.032806.104614>.
- M. R. So/rensen and A. F. Voter. Temperature-accelerated dynamics for simulation of infrequent events. *J. Chem. Phys.*, 112(21):9599–9606, May 2000. URL <https://doi.org/10.1063/1.481576>.
- Y. Sugita and Y. Okamoto. Replica-exchange molecular dynamics method for protein folding. *Chem. Phys. Lett.*, 314:141–151, 1999. URL <http://adsabs.harvard.edu/abs/1999CPL...314..141S>.
- W. Teschner, R. Rudolph, and J. R. Garel. Intermediates on the folding pathway of octopine dehydrogenase from *Pecten jacobaeus*. *Biochemistry*, 26(10):2791–2796, May 1987. URL <https://doi.org/10.1021/bi00384a021>.
- Guangyao Zhou**, Z. Lu, and Y. Peng. L1-graph Construction Using Structured Sparsity. *Neurocomputing*, 120:441–452, 2013.
- Guangyao Zhou**, S. Geman, and J. M. Buhmann. Sparse Feature Selection by Information Theory. *IEEE International Symposium on Information Theory (ISIT)*, 2014, pages 926–930, 2014. doi: 10.1109/ISIT.2014.6874968.
- Guangyao Zhou**, E. Borenstein, O. Livne, and A. Brandt. Multiscale Optimization for Stochastic Ill-conditioning in Deep Neural Networks. *In preparation*, 2018a.
- Guangyao Zhou**, M. Harrison, and S. Geman. Differentiable Inference and Compositionality in Probabilistic Generative Models. *In progress*, 2018b.
- Guangyao Zhou**, J. Loper, Matthew Harrison, and S. Geman. Base-pair Ambiguity and the Kinetics of RNA Folding. *bioRxiv*, 2018c. doi: 10.1101/329698. URL <https://www.biorxiv.org/content/early/2018/05/25/329698>.

- T. S. van Erp and P. G. Bolhuis. Elaborating transition interface sampling methods. *J. Comput. Phys.*, 205(1):157–181, May 2005. URL <http://www.sciencedirect.com/science/article/pii/S0021999104004620>.
- A. F. Voter. Hyperdynamics: Accelerated Molecular Dynamics of Infrequent Events. *Phys. Rev. Lett.*, 78(20):3908–3911, May 1997. URL <https://link.aps.org/doi/10.1103/PhysRevLett.78.3908>.
- A. F. Voter. Parallel replica method for dynamics of infrequent events. *Phys. Rev. B Condens. Matter*, 57(22):R13985–R13988, June 1998. URL <https://link.aps.org/doi/10.1103/PhysRevB.57.R13985>.
- D. J. Wales. Energy landscapes: calculating pathways and rates. *Int. Rev. Phys. Chem.*, 25(1-2):237–282, Jan. 2006. URL <https://doi.org/10.1080/01442350600676921>.
- A. Warshel and M. Levitt. Theoretical studies of enzymic reactions: dielectric, electrostatic and steric stabilization of the carbonium ion in the reaction of lysozyme. *J. Mol. Biol.*, 103(2):227–249, May 1976. URL <https://www.ncbi.nlm.nih.gov/pubmed/985660>.
- J. G. Wendel. Hitting Spheres with Brownian Motion. *Ann. Probab.*, 8(1):164–169, 1980. URL <http://www.jstor.org/stable/2243068>.
- A. M. A. West, R. Elber, and D. Shalloway. Extending molecular dynamics time scales with milestoning: Example of complex kinetics in a solvated peptide. *J. Chem. Phys.*, 126(14):145104, Apr. 2007. URL <https://doi.org/10.1063/1.2716389>.
- E. P. Wigner. The Transition State Method. In A. S. Wightman, editor, *Part I: Physical Chemistry. Part II: Solid State Physics*, pages 163–175. Springer Berlin Heidelberg, Berlin, Heidelberg, 1997. URL https://doi.org/10.1007/978-3-642-59033-7_16.
- L. T. Wille. Intractable computations without local minima. *Phys. Rev. Lett.*, 59(3):372, July 1987. URL <http://dx.doi.org/10.1103/PhysRevLett.59.372>.
- S. Yang, J. N. Onuchic, A. E. García, and H. Levine. Folding time predictions from all-atom replica exchange simulations. *J. Mol. Biol.*, 372(3):756–763, Sept. 2007. URL <http://dx.doi.org/10.1016/j.jmb.2007.07.010>.
- G. Zemora and C. Waldsich. RNA folding in living cells. *RNA Biol.*, 7(6):634–641, Nov. 2010. URL <https://www.ncbi.nlm.nih.gov/pubmed/21045541>.
- W. Zheng, M. Andrec, E. Gallicchio, and R. M. Levy. Recovering Kinetics from a Simplified Protein Folding Model Using Replica Exchange Simulations: A Kinetic Network and Effective Stochastic Dynamics. *J. Phys. Chem. B*, 113(34):11702–11709, Aug. 2009. URL <https://doi.org/10.1021/jp900445t>.
- M. Zuker, D. H. Mathews, and D. H. Turner. Algorithms and Thermodynamics for RNA Secondary Structure Prediction: A Practical Guide. In J. Barciszewski and B. F. C. Clark, editors, *RNA Biochemistry and Biotechnology*, NATO Science Series, pages 11–43. Springer Netherlands, 1999. URL http://link.springer.com/chapter/10.1007/978-94-011-4485-8_2.