

Preface

Overview

The evolutionary history of a set of genes, species, or individuals provides a context in which biological questions can be addressed. For this reason, phylogeny estimation is a fundamental step in many biological studies, with many applications throughout biology, such as protein structure and function prediction, analyses of microbiomes, inference of human migrations, etc. In fact, there is a famous saying by Dobzhansky that “Nothing in biology makes sense except in the light of evolution” (Dobzhansky, 1973).

Because phylogenies represent what has happened in the past, they cannot be directly observed but rather must be estimated. Consequently, increasingly sophisticated statistical models of sequence evolution have been developed, and are now used to estimate phylogenetic trees. Indeed, over the last few decades, hundreds of software packages and novel algorithms have been developed for phylogeny estimation, and this influx of computational approaches into phylogenetic estimation has transformed systematics. The availability of sophisticated computational methods, fast computers and high-performance computing (HPC) platforms, and large sequence datasets enabled through DNA sequencing technologies, has led to the expectation that highly accurate large-scale phylogeny estimation, potentially answering open questions about how life evolved on earth, should be achievable.

Yet large-scale phylogeny estimation turns out to be much more difficult than expected, for multiple reasons. First, all the best methods are computationally intensive, and standard techniques do not scale well to large datasets; for example, maximum likelihood phylogeny estimation is NP-hard, so exact solutions cannot be found efficiently (unless $P = NP$), and Bayesian MCMC methods can take a long time to reach stationarity. While massive parallelism can ameliorate these challenges to some extent, it doesn't really address the basic challenge inherent in searching an exponential search space. However, another issue is that the statistical models of sequence evolution that properly address genomic data are substantially more complex than the ones that model individual loci, and methods to estimate genome-scale phylogenies are (relatively speaking) in their infancy compared to methods for single gene phylogenies. Finally, there is a substantial gap between performance as suggested by mathematical theory (which is used to establish guarantees about

methods under statistical models of evolution) and how well the methods actually perform on data – even on data generated under the same statistical models! Indeed, this gap is one of the most interesting things about doing research in computational phylogenetics, because it means that the most impactful research in the area must draw on mathematical theory (especially probability theory and graph theory) as well as on observations from data.

The main goal of this text is to enable researchers (typically graduate students in computer science, applied mathematics, or statistics) to be able to contribute new methods for phylogeny estimation, and in particular to develop methods that are capable of providing improved accuracy for large heterogeneous datasets that are characteristic of the types of inputs that are increasingly of interest in practice. The secondary goal is to enable biologists to understand the methods and their statistical guarantees under these models of evolution, so that they can select appropriate methods for their datasets, and select appropriate datasets given the available methods.

Some of the material in the textbook is fairly mathematical, and presumes undergraduate coursework in discrete mathematics and algorithm design and analysis. However, no background in biology is assumed, and the assumed statistics background is relatively lightweight. While some students without the expected background in computer science may find it difficult to understand some of the proofs, my goal has been to enable all students to understand the theoretical guarantees for phylogeny estimation methods and the statistical models on which they are based, so that they can adequately critique the scientific literature, and also choose methods and datasets that are best able to address the scientific questions they wish to answer.

Outline of the Textbook

Part I provides the “Discrete Mathematics for Phylogenetics” foundations for the textbook; the concepts and mathematics introduced in this part are the building blocks for algorithm design in phylogenetics, especially for developing methods that can scale to large datasets; understanding these concepts makes it possible to understand theoretical guarantees of methods under statistical models of evolution. Chapter 1 introduces the major themes involved in computational phylogenetics, addressing both theory (e.g., statistical consistency under a statistical model of evolution) and performance on both simulated and biological data. This chapter uses the Cavender–Farris–Neyman model of binary sequence evolution since understanding issues in analyzing data generated by this very simple model is helpful to understanding statistical estimation under the commonly used models of molecular sequence evolution. Chapter 2 introduces trees as graph-theoretic objects, and presents different representations of trees that will be useful for method development. Chapters 3, 4, and 5 present different types of methods for phylogenetic tree estimation (based on combining subtrees, using character data, or using distances, respectively). Chapter 6 presents methods for analyzing sets of trees, each on the same set of taxa, and for computing consensus trees and agreement subtrees; it also discusses how these methods are used to estimate support for different phylogenetic hypotheses. Chapter 7 examines the

topic of supertree estimation, where the input is a set of trees on overlapping sets of taxa and the objective is a tree on the full set of taxa. Supertree methods are of interest in their own right and also because they are key algorithmic ingredients in divide-and-conquer methods, a topic we return to in Chapter 11.

Part II of the textbook is concerned with molecular phylogenetics. Chapter 8 presents commonly used statistical models of molecular sequence evolution and statistical methods for phylogeny estimation under these models. However, standard sequence evolution models do not include events such as insertions, deletions, and duplications, which can change the sequence length. These are very common processes, so biological sequences are usually of different lengths and must first be put into a *multiple sequence alignment* before they can be analyzed using phylogeny estimation methods; the subject of how to compute a multiple sequence alignment is covered in Chapter 9. Constructing a species tree or even a phylogenetic network from different gene trees in the presence of gene tree heterogeneity due to incomplete lineage sorting, gene duplication and loss, horizontal gene transfer, or hybridization is a fascinating research area that we present in Chapter 10. We end with Chapter 11, which addresses method development for estimating trees on large datasets. Large-scale phylogeny estimation is increasingly important since nearly all good approaches to phylogeny and multiple sequence alignment estimation are computationally intensive (either heuristics for NP-hard optimization problems or Bayesian methods), and many large datasets are being assembled that cannot be accurately analyzed using existing methods.

Each chapter ends with a set of review questions and homework problems. The review questions are easy to answer and do not require any significant problem solving or calculation. The homework problems are largely pen and paper problems that reinforce the mathematical content of the text.

The textbook comes with four appendices. Appendix A provides an introduction to biological evolution and data; the textbook can be read without it, but the reader who wishes to analyze biological data will benefit from this material. Appendix B provides an introduction to algorithm design and analysis; this material is not necessary for students with undergraduate computer science backgrounds, but may be a helpful introduction for students without this background. Appendix C provides some guidelines about how to write papers that introduce new methods or evaluate existing methods. Appendix D provides computational projects ranging from short term (i.e., a few days) to research projects that could lead to publications. In fact, several of the final projects for my Computational Phylogenetics courses have grown into journal publications (e.g., Bayzid et al. (2014); Zimmermann et al. (2014); Davidson et al. (2015); Chou et al. (2015); Nute and Warnow (2016)).

I wish to thank my editor, David Tranah, for his detailed and insightful comments on the many earlier versions of the text. I also wish to thank my students, colleagues, and family members who gave helpful criticism, including Edward Braun, Sarah Christensen, Steve Evans, Dan Gusfield, Joseph Herman, Ally Kaminsky, Laura Kubatko, Ari Löytynoja, Siavash Mirarab, Erin Molloy, Luay Nakhleh, Mike Nute, David Posada, Bhanu Renukuntla, Ehsan Saleh, Erfan Sayyari, Kimmen Sjölander, Travis Wheeler, and Mark Wilkinson. The several anonymous reviewers also gave very useful comments.

The images of the Monterey Cypress tree on the front and back covers are in honor of the CIPRES project (www.phylo.org), an NSF-funded project for phylogenetic research that I co-led with Bernard Moret from 2003–2010. Many of the algorithmic advances discussed in the text came out of research supported by CIPRES.