

Preface

With a large number of prokaryotic and eukaryotic genomes completely sequenced and more forthcoming, access to the genomic information and synthesizing it for the discovery of new knowledge have become central themes of modern biological research. Mining the genomic information requires the use of sophisticated computational tools. It therefore becomes imperative for the new generation of biologists to be familiar with many bioinformatics programs and databases to tackle the new challenges in the genomic era. To meet this goal, institutions in the United States and around the world are now offering graduate and undergraduate students bioinformatics-related courses to introduce them to relevant computational tools necessary for the genomic research. To support this important task, this text was written to provide comprehensive coverage on the state-of-the-art of bioinformatics in a clear and concise manner.

The idea of writing a bioinformatics textbook originated from my experience of teaching bioinformatics at Texas A&M University. I needed a text that was comprehensive enough to cover all major aspects in the field, technical enough for a college-level course, and sufficiently up to date to include most current algorithms while at the same time being logical and easy to understand. The lack of such a comprehensive text at that time motivated me to write extensive lecture notes that attempted to alleviate the problem. The notes turned out to be very popular among the students and were in great demand from those who did not even take the class. To benefit a larger audience, I decided to assemble my lecture notes, as well as my experience and interpretation of bioinformatics, into a book.

This book is aimed at graduate and undergraduate students in biology, or any practicing molecular biologist, who has no background in computer algorithms but wishes to understand the fundamental principles of bioinformatics and use this knowledge to tackle his or her own research problems. It covers major databases and software programs for genomic data analysis, with an emphasis on the theoretical basis and practical applications of these computational tools. By reading this book, the reader will become familiar with various computational possibilities for modern molecular biological research and also become aware of the strengths and weaknesses of each of the software tools.

The reader is assumed to have a basic understanding of molecular biology and biochemistry. Therefore, many biological terms, such as *nucleic acids*, *amino acids*, *genes*, *transcription*, and *translation*, are used without further explanation. One exception is *protein structure*, for which a chapter about fundamental concepts is included so that

algorithms and rationales for protein structural bioinformatics can be better understood. Prior knowledge of advanced statistics, probability theories, and calculus is of course preferable but not essential.

This book is organized into six sections: biological databases, sequence alignment, genes and promoter prediction, molecular phylogenetics, structural bioinformatics, and genomics and proteomics. There are nineteen chapters in total, each of which is relatively independent. When information from one chapter is needed for understanding another, cross-references are provided. Each chapter includes definitions and key concepts as well as solutions to related computational problems. Occasionally there are boxes that show worked examples for certain types of calculations. Since this book is primarily for molecular biologists, very few mathematical formulas are used. A small number of carefully chosen formulas are used where they are absolutely necessary to understand a particular concept. The background discussion of a computational problem is often followed by an introduction to related computer programs that are available online. A summary is also provided at the end of each chapter.

Most of the programs described in this book are online tools that are freely available and do not require special expertise to use them. Most of them are rather straightforward to use in that the user only needs to supply sequences or structures as input, and the results are returned automatically. In many cases, knowing which programs are available for which purposes is sufficient, though occasionally skills of interpreting the results are needed. However, in a number of instances, knowing the names of the programs and their applications is only half the journey. The user also has to make special efforts to learn the intricacies of using the programs. These programs are considered to be on the other extreme of user-friendliness. However, it would be impractical for this book to try to be a computer manual for every available software program. That is not my goal in writing the book. Nonetheless, having realized the difficulties of beginners who are often unaware of or, more precisely, intimidated by the numerous software programs available, I have designed a number of practical Web exercises with detailed step-by-step procedures that aim to serve as examples of the correct use of a combined set of bioinformatics tools for solving a particular problem. The exercises were originally written for use on a UNIX workstation. However, they can be used, with slight modifications, on any operating systems with Internet access.

In the course of preparing this book, I consulted numerous original articles and books related to certain topics of bioinformatics. I apologize for not being able to acknowledge all of these sources because of space limitations in such an introductory text. However, a small number of articles (mainly recent review articles) and books related to the topics of each chapter are listed as "Further Reading" for those who wish to seek more specialized information on the topics. Regarding the inclusion of computational programs, there are often a large number of programs available for a particular task. I apologize for any personal bias in the selection of the software programs in the book.

One of the challenges in writing this text was to cover sufficient technical background of computational methods without extensive display of mathematical formulas. I strived to maintain a balance between explaining algorithms and not getting into too much mathematical detail, which may be intimidating for beginning students and nonexperts in computational biology. This sometimes proved to be a tough balance for me because I risk either sacrificing some of the original content or losing the reader. To alleviate this problem, I chose in many instances to use graphics instead of formulas to illustrate a concept and to aid understanding.

I would like to thank the Department of Biology at Texas A&M University for the opportunity of letting me teach a bioinformatics class, which is what made this book possible. I thank all my friends and colleagues in the Department of Biology and the Department of Biochemistry for their friendship. Some of my colleagues were kind enough to let me participate in their research projects, which provided me with diverse research problems with which I could hone my bioinformatics analysis skills. I am especially grateful to Lisa Peres of the Molecular Simulation Laboratory at Texas A&M, who was instrumental in helping me set up and run the laboratory section of my bioinformatics course. I am also indebted to my former postdoctoral mentor, Carl Bauer of Indiana University, who gave me the wonderful opportunity to learn evolution and phylogenetics in great depth, which essentially launched my career in bioinformatics. Also importantly, I would like to thank Katrina Halliday, my editor at Cambridge University Press, for accepting the manuscript and providing numerous suggestions for polishing the early draft. It was a great pleasure working with her. Thanks also go to Cindy Fullerton and Marielle Poss for their diligent efforts in overseeing the copyediting of the book to ensure a quality final product.

Jin Xiong