

Preface

This is volume 1 of our new book on the applied hierarchical modeling of the three central quantities in ecology—abundance, or density, occurrence, and species richness—as well as of parameters governing their change over time, especially survival and recruitment. Hierarchical modeling is a growth industry in ecology. In the last 10 years there have been a dozen or more books focused on hierarchical modeling in ecology including Banerjee et al. (2004), Clark and Gelfand (2006), Clark (2007), Gelman and Hill (2007), McCarthy (2007), Royle and Dorazio (2008), King et al. (2009), Link and Barker (2010), Kéry and Schaub (2012), Hobbs and Hooten (2015), etc. How can we possibly add 700+ more pages (and perhaps 1500 if you count volume 2) to what is known on this topic? That's a good question!

In this book we cover several classes of models that have previously received only cursory or no treatment at all in the hierarchical modeling literature, and yet they are extremely important in ecology (e.g., distance sampling). Moreover, we give complete recipes for analyzing these models and all others covered in the book, using program R in general and the R package *unmarked* in particular, and very extensively using the generic Bayesian modeling software BUGS. The use of *unmarked* is completely novel compared to these other books. Some models that we cover here and especially in volume 2 were simply unimaginable a couple years ago, e.g., the open hierarchical distance sampling models, the metacommunity abundance models and models with explicit population dynamics (i.e., the famous model of Dail and Madsen, 2011), which can be fit to counts and related data, including even distance sampling data. Much of this material is extremely new, and some has only just appeared in the literature in the last year or so. Thus, this book represents a timely synthesis and extension of the state of hierarchical modeling in ecology that builds on previous efforts, but covers much new and important territory, and provides implementations using both the likelihood (*unmarked*) and Bayesian (BUGS) frameworks.

A BOOK OF MONOGRAPHS

In a sense, *Applied Hierarchical Modeling for Ecologists (AHM)* is a book of books, or a *book of monographs*. Volume 1 contains the first two parts, a prelude, which introduces the necessary concepts and techniques in five chapters, followed by six chapters that deal with static demographic models of distribution, abundance, and species richness and other descriptors of communities and metacommunities. Volume 2 will contain two further parts on dynamic models and on advanced demographic models for populations and communities; see below for more information on the division of topics between volumes 1 and 2 of *AHM* and on the content that we envision for volume 2.

Looking back at volume 1 now, at the time of writing of this preface, we feel as if we have packaged almost a dozen independent books into this one book. There are general, introductory “monographs” on the concepts of distribution, abundance, and species richness and their measurement and modeling in practice (Chapter 1), on hierarchical models and their analysis (Chapter 2), on linear, generalized linear, and mixed models (Chapter 3), on data simulation in R (Chapter 4), and on the celebrated BUGS language and software (Chapter 5). After that, there are six comprehensive monographs on demographic models for distribution, abundance, and species richness in the context of what we call a “meta-population design,” that is, the extremely common situation where you measure something in a population or a community at more than a single point in space.

In the prelude, and following the introductory Chapter 1, we have one monograph to cover hierarchical models (HMs) and their Bayesian and frequentist analyses (Chapter 2). The next

monograph (Chapter 3) provides a highly accessible review of that “heart” of applied statistics: linear models, generalized linear models (GLMs), and simple mixed models, all of them illustrated in the context of one extremely simple ecological data set. Data simulation is one of the defining features of this book because it provides such immense benefits for the work of ecologists (and also for statisticians). Hence, the next monograph (Chapter 4) is dedicated to this essential topic and walks you through the R code necessary for the generation of one simple type of data set that is fundamental to the classes of models covered in this book: the case where one goes out and counts birds (or any other species) at multiple places (e.g., 20, 100, or 267) and repeats these counts at each site multiple times (e.g., 2 or 3).

The BUGS model definition language is implemented in three currently used BUGS engines for Bayesian inference: WinBUGS (Lunn et al., 2000), OpenBUGS (Thomas et al., 2006), and JAGS (Plummer, 2003). It has also just been adopted in the exciting new R package NIMBLE (NIMBLE Development Team, 2015; de Valpine et al., in review), which is a general model fitting software that uses and extends the BUGS language for flexible specification of HMs and allows analysis of HMs with both maximum likelihood and Bayesian posterior inference.

Over the first 25 years of its existence, BUGS has been instrumental in the surge of Bayesian statistics in all kinds of sciences including ecology (Lunn et al., 2009). It has grown by far into the most important general, Bayesian modeling language, and its user population keeps growing at a rapid rate (and of course we hope to increase that rate even more with this book). BUGS is unique in giving you as a nonstatistician a modeling freedom that lets you develop, test, and fit models that you wouldn’t even have dared to dream of in the pre-BUGS era of ecological modeling (which we might call the ecological Stone Age;...). Although there are now many useful introductory books on BUGS (e.g., McCarthy, 2007; Kéry, 2010; Lunn et al., 2013; Korner-Nievergelt et al., 2015), we have decided to write yet another practical BUGS introduction and package it into Chapter 5. It is our latest and best attempt at covering as much as possible on this topic and including some of the latest tricks in BUGS modeling in a mere 70 book pages, illustrating the use of all three BUGS engines and focusing on the models covered in Chapter 3, i.e., linear models, GLMs, and simple mixed models. By introducing BUGS for exactly the kinds of models that you are likely to be familiar with already, we hope to make it especially easy for you to grasp the Bayesian side of the analysis and the implementation of these essential models in the BUGS language.

In the second part of the book, we present six monographs that contain a comprehensive treatment of important classes of models for inference about distribution, abundance, and species richness, and related demographic population or community metrics in so-called “meta-population designs” (Royle, 2004a; Kéry and Royle, 2010), i.e., for the frequent case where you are interested in these things not at a single place but have studied them at multiple sites. Specifically, in Chapter 6 we cover binomial mixture, or *N*-mixture, models (Royle, 2004b), which are a unique type of model for count data on unmarked individuals (that is, you do not need to keep track of which individual is which across the repeated measurements of abundance at a site) and that contains an explicit measurement error model, which corrects your inferences for the biases that would otherwise be caused by undercounting due to imperfect detection probability. Chapter 7 covers a “sister-type” of model, the multinomial mixture model (Royle, 2004a; Dorazio et al., 2005), which only differs from binomial mixture models in the type of data to which it is fitted: typically you need individual recognition, that is, you have capture-recapture-type of data, but again collected not at a single site but at multiple places. Both types of mixture models have been around for about 10 years now and previously they have been featured in

some of the above-cited hierarchical modeling books (mostly in Royle and Dorazio, 2008), but never before have they been covered in such detail and, especially, in a manner that makes them so accessible to you as an ecologist.

Chapters 8 and 9 are special in that they provide perhaps the first large, and yet practical and applied, synthesis in distance sampling more than 10 years after the two classics by Buckland et al. (2001, 2004a) were published. In our two distance sampling monographs, we provide a fresh, new look at distance sampling in the context of hierarchical models in an essentially book-length treatment. We hope that this will help to make this important type of model even more widely understood and used by ecologists. While we cover mainly static models in volume 1 and then cover dynamic models in more detail in volume 2 of *AHM*, we have deviated slightly from this rule in Chapters 8 and 9, where we have preferred topical unity over conceptual unity by keeping all of distance sampling (closed and open) together. Nevertheless, we plan to cover several more cutting-edge open and other novel extensions of hierarchical distance sampling (HDS) models in volume 2.

These two monographs, and more specifically the wealth of material on hierarchical distance sampling, are perhaps those with most novelty in our book. Though again invented just over 10 years ago (around 2004), HDS has recently experienced a boost with the widespread realization that this type of specification of distance sampling models enables extremely flexible modeling of spatially or temporally replicated distance sampling data or combined analyses of data sets collected under differing protocols ("integrated models"), which was thought impossible before or at least was never achieved. For instance, it is perfectly doable (or even trivial) to model population dynamics (Sollmann et al., 2015) or community size and composition (Sollmann et al., in press) from distance sampling data within the context of hierarchical models. And, the power of BUGS nowadays makes the implementation of such models possible even for ecologists, since really such models differ in only relatively minor ways from similar models for other data types (e.g., of the capture-recapture type).

Hence, we hope that we contribute to change your view of "capture-recapture" and "distance sampling" as being two widely separated fields to a new way of seeing them as really relatively minor variations on the overarching theme of hierarchical models, which have one model component for abundance, or density, and in another model component describe the measurement error that induces imperfect detection and therefore undercounting (Borchers et al., 2015). The only thing that changes when you move from a capture-recapture to a distance sampling model is the specific parameterization of the measurement error underlying the observed data and of course the type of data that you need to estimate the parameters of that measurement error model. This wonderful, unifying power of describing statistical models in a hierarchical way lets you much better grasp the similarities among large numbers of models that were often thought as totally distinct hitherto. It is one of the main themes of this book and one on which we will say much more throughout the book. For instance, we hope that you will recognize that there are really only quite minor differences between a binomial mixture model for counts of unmarked individuals, a multinomial mixture model for capture-recapture data, and a hierarchical distance sampling model—the only difference is again the measurement error model, while the state model, that is, the description of the essential biological quantity (abundance or density), is exactly the same in all three types of models.

The penultimate monograph (Chapter 10) is on occupancy modeling (MacKenzie et al., 2002; Tyre et al., 2003). This powerful type of model for occurrence or distribution comes with an explicit measurement error component model for both false-negatives and false-positives (models by Royle and Link, 2006; Aing et al., 2011; Miller et al., 2011, 2013b; Sutherland et al., 2013; Chambert et al.,

2015) or with a measurement error model for false negatives only (all other types of occupancy models). Occupancy models have become huge in ecology and have experienced a steep growth curve in both the number of papers that further develop the theory of these models and especially also in studies that apply this design and the associated models. (We have even heard rumors that the vigorous growth of the field has “scared” some ecology journal editors so that they put a cap on the number of occupancy papers they accept—a strange way of stifling progress one would think.) Occupancy models have received one book-length treatise so far (MacKenzie et al., 2006), with a second edition that is in preparation, and several customized software products that specialize in them, especially PRESENCE (Hines, 2006) and MARK (White and Burnham, 1999; Cooch and White, 2014). In this first *AHM* volume, we deal with single-species occupancy models in great detail and cover some topics (e.g., some models for data collected along space or time “transects”) that haven’t been covered in any book before. In volume 2, we will add several more monographs on a large variety of occupancy model types; see below.

The final monograph in volume 1 covers community models, that is, community or multispecies variants of all the previous models. Specifically, we cover the community variant of an occupancy model (Chapter 10) and the community variant of a binomial N -mixture model (Chapter 6). These powerful hierarchical models enable inferences at multiple scales, that of the individual species, that of a local community, and that of an entire metacommunity. As always in this book, both come with an explicit measurement error model for the desired state of inference, presence/absence, or abundance of each individual species at every site in the “meta-population.” These models have experienced much increased attention in the very recent past (Iknayan et al., 2014; Yamaura et al., 2012, in press), and we provide a much needed, comprehensive and yet supremely practical monograph on both the abundance and on the occupancy-based community models.

Of course, apart from serving as a standalone introduction to this large range of powerful and useful hierarchical models, the material in volume 1 also lays the groundwork for more models and more advanced material in volume 2. See below for more about the division of content between the two volumes.

UNIFYING THEMES

AHM is not just a hodgepodge of models that have not previously been covered in detail or at all. Rather, our development and organization of these models has a number of unifying themes that we emphasize throughout the book:

- hierarchical modeling
- data simulation
- measurement error models
- dual inference paradigm approach (Bayesianism and frequentism)
- accessible and gentle style (including hierarchical likelihood construction and data simulation)
- “cookbook recipes”
- predictions

One, we advocate *hierarchical modeling* as a unifying concept and overarching principle in modeling and also conceptually; as we have emphasized before, when seen as HMs all these models almost look the same (or very similar) and it is quite trivial to move from one to another, e.g., from a capture-recapture model to a distance sampling model to an occupancy model or even to a community

or metacommunity model. We dedicate an entire chapter to introduce and explain the crucial concept of HMs, which permeates every section of this book.

Two, we use *data simulation* throughout the book, because this is so tremendously important in practice, for statisticians, but much more so still for ecologists. This is done in hardly any other book we know of, except for two of our earlier books (Kéry, 2010; Kéry and Schaub, 2012). Though quite frequently done by statisticians and also by ecologists in many different modes, we believe that data simulation should be done *much* more widely still. We dedicate an entire chapter to data simulation (Chapter 4) and therein explain the major advantages for you when you start doing this routinely for your work. Among them, perhaps the two most important benefits of data simulation are, first, that it enforces on you a complete understanding of your model. If you don't understand your model, you will not be able to write R code to simulate data under that model—it's as simple as that. We would even go as far as saying that a data simulation algorithm provides a complete description of a statistical model. Indeed, throughout the book we use data simulation in R in a completely novel fashion *to explain a statistical model!*

The second major benefit of data simulation is that it serves an important role to validate both your MCMC algorithm (whether written by yourself or produced by an MCMC black box such as BUGS) and to validate your model code. For most model classes in the book we provide R functions to simulate data under various types of models. We hope that these will be widely used in the many different modes of data simulation (as per Chapter 4)...or perhaps sometimes simply to marvel at the pretty and highly variable graphical output they produce.

Three, one defining feature of all main classes of models in our book is the presence of a submodel that contains an explicit description of the *measurement error* process underlying all data on the distribution and abundance of individual species and even more perhaps when you study them in entire communities or metacommunities. Unlike the types of measurement error for continuous variables (such as body length) to which you may have been exposed, the measurement error for discrete measurements (e.g., counts and presence/absence measurements) is of a radically different nature and comes in exactly two types: false-positive and false-negative measurement error, with the complement of the latter typically being called detection or encounter probability, or detectability for short. These are very different types of measurement error, which you cannot expect to cancel out in the mean over several measurements. Hence, unless you account for them in your models for distribution, abundance, and species richness, badly biased inferences may result. Our *AHM* book is an "estimationist book" in line with a rapidly increasing number of previous works that emphasize the measurement error processes in ecological models for distribution and abundance in ecology, such as Otis et al. (1978), Seber (1982), Buckland et al. (2001), Borchers et al. (2002), Williams et al. (2002), Buckland et al. (2004a), Amstrup et al. (2005), MacKenzie et al. (2006), Royle and Dorazio (2008), King et al. (2009), Kéry and Schaub (2012), McCrea and Morgan (2014), and Royle et al. (2014).

Four, we are neither purebred Bayesians nor hardcore frequentists, rather, we are big fans of a *dual inference paradigm approach*, i.e., the use of Bayesianism *and* frequentism alongside, as it seems especially useful for the particular case. While perhaps both of us have a slight personal slant toward Bayesianism, there are advantages and disadvantages of both Bayesianism and frequentism, and these may come into play more or less for any given data set or scientific question (Little, 2006; de Valpine, 2009, 2011). In addition, the choice of whether a Bayesian or a frequentist analysis is most appropriate will also be affected by the availability of a well-trained analyst and/or a fast computer, with Bayesian solutions often requiring more statistical and programming experience and faster computers. Thus, we

are firm believers in the value of a dual inference paradigm approach, and this is a pervasive theme of our book as well. The dual inference paradigm approach appears in all but one chapters of this book. This approach has been done a little bit in some previous books (especially in Royle and Dorazio, 2008) but never to our knowledge in such a completely integrated way as in this book. Every topical chapter in Part 2 covers a class of models using both inference paradigms and emphasizes things that are easier or harder to do one way or the other (the exception being Chapter 11, where it is very hard to do a non-Bayesian analysis of these parameter-rich models).

Five, we have striven to make this book *gentle and accessible in style and easy to read*. This means that we do of course present formulae and equations, but perhaps fewer than in many other comparable statistics books. Many ecologists cannot read even moderately complex likelihood expressions. This is perhaps not a good state of affairs, but it is a simple fact of life that is unlikely to change anytime soon. We believe that the hierarchical construction of the likelihood, as a series of conditional probability statements as in every topical chapter in this book (and as we naturally do when specifying these models in the BUGS language), is perhaps the *only* way in which a fairly large proportion of ecologists have any chance of being able to read and understand the likelihood of a somewhat complex model. In addition to algebra, we use especially data simulation (and hence R code) to describe our models throughout the book. We find that R code for data simulation is an extremely clear and transparent way of implicitly describing the likelihood of a model. This seems to be a completely novel idea that has never been expressed explicitly before.

Six, we illustrate analyses of each class of models using a complete set of steps that you would use in your work ("*cookbook recipes*"). This includes not just fitting the models but producing summary analyses such as response curves and prediction, and especially illustrating maps of abundance and occurrence, and also assessing the goodness-of-fit of models. We believe that providing cookbook recipes is frowned upon by many statisticians because there is a feeling that this encourages people to do things that they don't understand. We are convinced that this sentiment is mostly unfounded. First, and most importantly, for any but an extremely trivial analysis, the practitioner will still have to understand the model and the analysis in order to not make any of a myriad of trivial errors that will make the BUGS program crash. Second, without at least some understanding, he will probably not be able to describe the results in an intelligible way in the results section of his paper or to explain them to her supervisor, advisor, or colleague. On the other hand, even some of the most basic of statistical analyses, namely linear models with factor levels, are extremely widely misunderstood, i.e., people don't understand what the intercept means and what the treatment contrast parameters are. Hence, some abuse is to be expected with *any* kind of statistical model for which easy-to-use code is made widely available. In addition, in complex models in this and similar books, sometimes even very basic steps such as formatting the data into a three- or four-dimensional array can be a complete stumbling block to an R novice, even though he may have a decent conceptual grasp of a model. In this case, the availability of cookbook analysis code is essential. Finally, it is the experience of at least one of the authors that only fitting a model and looking at the estimates may sometimes really let one understand what these parameters mean. Of course, this latter effect may be magnified still when you fit the model to simulated data, where you know what you input into your data set and therefore what ballpark estimates you can expect. In summary, we believe that it is not evil to hand out cookbook recipes but rather that they *ought* to be given much more widely, and we do exactly this throughout our book.

Seven, and finally, one of the examples of us giving ample code recipes is for *prediction*, i.e., for the computation of the expected value of some quantity (e.g., the response or some parameter) for a

range of values for one or more covariates. Forming such predictions is extremely important for you in two ways: to even understand what the model is telling you about the form of some covariate relationship when you have log, logistic, or similar link functions, polynomial terms, or interactions; and second to present the results of your analysis, e.g., in a figure in your paper. We emphasize prediction throughout the book, especially predictions in geographic space, leading to maps of species abundance and occurrence (the associated models are then called “species distribution models”); this is a very hot topic nowadays. The forming of predictions, and how to put these predictions on a map, is the focus of every single monograph in the second part of the book and also appears extensively in the prelude chapters.

Although this book is especially geared toward ecologists, it presents the cutting edge of the current state and understanding of all of the models presented. At several places, we were not shy to lay open our partial lack of understanding about some topics, in the hope to emphasize the need for further research; this includes goodness-of-fit in these models (and probably in many other classes of hierarchical models in general), the good fit versus bad prediction dilemma with some negative binomial N -mixture models in Chapter 6, or the use of spatial instead of temporal replication for obtaining information about measurement error in occupancy models in Chapter 10. Clearly, our aim in writing this book is not to show off how much we know but to help you to learn these models to understand and apply them. This includes a recognition of where their limits or the general limits of our understanding about them are and where you could therefore make a contribution to the progress in this field.

THE unmarked PACKAGE

Somebody once said that he (or she) did not trust any R package unless it has a book written about it. So now you can finally trust the R package `unmarked` (Fiske and Chandler, 2011) because this is also a book about `unmarked`. The `unmarked` package is fully general, and as part of the R programming environment, it allows you to embed your analyses seamlessly into your R programming. This is a great help when running simulations, for data processing and formatting, running analyses in batch mode (e.g., looping over many species, years, sites), documenting your data processing and analysis steps, and when analyzing results to produce plots, summary analyses, fit assessments, and model selection.

The `unmarked` package permits you to fit a large variety of closed and open hierarchical models and, to the best of our knowledge, it is the only package for likelihood estimation of (almost) all classes of models we cover in this book, although PRESENCE (Hines, 2006), MARK (White and Burnham, 1999; Cooch and White, 2014), and E-SURGE (Choquet et al., 2009b; Gimenez et al., 2014) fit occupancy models, and the former two also Royle-Nichols and binomial N -mixture models. One of the benefits of using `unmarked` for analyzing these various hierarchical models is that it streamlines and standardizes the work flow across models. An analysis of any class of hierarchical model in `unmarked` has a few basic steps, which include: (1) processing and packaging the data into an “`unmarked frame`” using standard constructor functions that ensure data are in the proper format; (2) utilization of a standard model fitting function that produces parameter estimates, standard errors, AIC, and other summary statistics; (3) summary analyses that include producing model selection tables, goodness-of-fit analyses (e.g., using parametric bootstrapping), and plotting predictions or fitted values. Each of these summary analyses is supported by standard functions that are part of the `unmarked` package.

The unmarked package is supported by an active and most of the time very friendly e-mail user group (groups.google.com/forum/#!forum/unmarked), which you can subscribe to for following developments and bug reports, or for requesting assistance. Finally, unmarked is an open source software development project. The source code is readily available and can be easily modified and extended by anyone. We encourage you to participate in the unmarked community.

COMPUTING¹

In a sense this is a book about ecological computing. While we emphasize the formulation and analysis of models, a vast majority of the effort to do so requires programming in the R language and running various functions in unmarked and in WinBUGS or JAGS. For Bayesian analysis we adopt the implementations of the BUGS language using WinBUGS and JAGS (and could equally well have used OpenBUGS or NIMBLE). These are used almost equivalently with the help of the R packages R2WinBUGS and jagsUI, and there are only a very small number of minor differences between the JAGS and WinBUGS implementations of the BUGS language (see the JAGS manual, available on the Internet, and Lunn et al., 2013).

Interestingly, the use of BUGS is often frowned upon by statisticians as some kind of inferior approach to things, as compared to writing your own MCMC algorithm, and reliance on BUGS is readily criticized in reviews of papers and conference presentations (especially those that are widely attended by statisticians). The academic statistician's view is often that you should be writing your own MCMC because then you understand what's going on under the hood. We disagree with this view. Now, and we think even 20 years into the future, the vast majority of ecologists will not be able to write their own MCMC nor even will most ecologists want to do that. Indeed, many statisticians can't do that either. On the other hand, BUGS makes accessible to ecologists the extremely convenient and useful technique of MCMC and consequently the ability to describe models and analyze them without having to have a PhD statistician helping them out. We therefore strongly advocate for the use of the BUGS language in whatever implementation is convenient (WinBUGS, JAGS, OpenBUGS, NIMBLE, or some future implementation). To be sure, custom MCMC algorithms may be *much* more efficient for any particular model or application. However, the time to produce custom algorithms really renders that approach impractical for most situations and for most people. Moreover, perhaps the greatest thing about the BUGS programs is the BUGS model definition language. This has proved to be supremely easy to understand for statisticians and nonstatisticians alike in their attempts to formulate, with confidence, even very complex statistical and simulation models. Therefore, we feel that the BUGS *language* is here to stay. There may be some inefficiency to the current implementations, but who's to say that a more efficient implementation won't be invented in the future? And, of course computing power is always improving and will continue to do so. In particular, computers will certainly have many more cores, and therefore multicore processing will improve the runtime of many models.

As to the diffuse fears of some when using a computational MCMC black box such as BUGS, we have argued before that data simulation has an important role to play in modern ecological modeling. Analyzing simulated data can give you much confidence about the good or bad behavior of a computational procedure—of course not for every single particular case (but you can't have

¹Use of product names does not imply endorsement by the US government.

this anyway, e.g., your likelihood maximization algorithm may always get stuck at a local maximum or along some flat ridge), but on average, and that is what really counts. For instance, over the years we have used BUGS to fit models to literally many thousands of simulated data sets for the types of models presented in this book. And we have only exceedingly rarely experienced cases where the algorithm converged to a place in parameter space that was not close to or right at the correct value, i.e., the value used to simulate the data set. Thus, we are not at all made nervous by the occasional claims heard about how terribly difficult it is to achieve chain convergence in an MCMC analysis.

ORGANIZATION

When we started *AHM* we did not think of it as comprising two volumes. But then we realized the wealth of material we had at our hands, and so now *AHM* comes in two volumes. As planned now, there are 25 chapters that are grouped in four parts, with two parts per volume (see Table 1). As already explained, the split of the whole *AHM* project into two volumes has introductory material including basic concepts of statistical modeling and inference and data simulation (Part I) and then single- and multispecies models of abundance and occurrence in static systems (Part II) in volume 1 (with the slight exception mentioned for distance sampling in Chapter 9). Volume 2 of the book will focus on dynamic and spatial models and other “advanced” topics.

WHO SHOULD READ THIS BOOK?

This book has two target audiences: first, ecologists and scientists and managers in related disciplines, where the demographic analysis of populations, meta-populations, communities, and metacommunities is a focus of interest. And second, statisticians, especially those hitherto unacquainted with these classes of HMs, which are hardly ever taught in standard methodology classes or in typical classical applied statistics texts. For the former group, the book represents a practical how-to guide for each class of models and thus should be accessible to anyone with basic R programming knowledge. Use of the BUGS language is needed also, but we hope you can gather the requisite skills by reading the earlier chapters of the book (3–5) or else you may consult an introductory BUGS book such as McCarthy (2007), Kéry (2010), Lunn et al. (2013), or Komer-Nievergelt et al. (2015). Because R programming is the standard now in many university curricula, we think the book should be ideal for a graduate level class on quantitative methods, either as a complete semester long course or part of a course covering specific models such as hierarchical modeling of abundance using *N*-mixture models, on occupancy models, and on hierarchical distance sampling.

BOOK WEB SITE AND USER GROUP E-MAIL LIST

For every analysis in the book we provide the complete instructions for organizing the data, fitting the model, and summarizing the results. Most of the commands are given directly in the book, although our companion Web site (<http://www.mbr-pwrc.usgs.gov/pubanalysis/keryroylebook/>) also provides the scripts for ready download. In addition, you find other information there, notably the solution to exercises, a list of errata as we find them (or, more likely, you detect and report them to us), etc.

Table 1 Outline of volume 1 and volume 2 of the Book *Applied Hierarchical Modeling in Ecology (AHM)*.

AHM volume 1: Prelude and static models

Preface

Part 1: Prelude

1. Distribution, abundance, and species richness in ecology
2. What are hierarchical models and how do we analyze them?
3. Linear models, generalized linear models (GLMs), and random effects: the components of hierarchical models
4. Introduction to data simulation
5. The Bayesian modeling software BUGS and JAGS

Part 2: Models for static systems

6. Modeling abundance using binomial N -mixture models
7. Modeling abundance using multinomial N -mixture models
8. Modeling abundance using hierarchical distance sampling
9. Advanced hierarchical distance sampling
10. Modeling distribution and occurrence using site-occupancy models
11. Community models

AHM volume 2: Dynamic and advanced models

Part 3: Models for dynamic systems

12. Modeling population dynamics with Poisson generalized linear mixed models (GLMMs) and some extensions
13. Modeling population dynamics with replicate counts within a season
14. Modeling population dynamics with distance sampling data
15. Hierarchical models of survival
16. Modeling species distribution and range dynamics using dynamic occupancy models
17. Modeling metacommunity dynamics using dynamic community models

Part 4: Advanced models

18. Multistate occupancy models
19. Modeling false-positives
20. Models for species interactions
21. Spatial models I
22. Spatial models II
23. Combination approaches/Integrated models
24. Spatial distance sampling and spatial capture-recapture
25. Conclusions

The *AHM* book (or indeed our larger “hierarchical modeling enterprise”) has an associated e-mail user group (<http://groups.google.com/forum/?hl=en#!forum/hmecology>), which you can subscribe to for following developments and bug reports, or for requesting assistance. There is some overlap with the unmarked e-mail user group, but the hierarchical modeling user group is more general and, in particular, is the only one specifically for questions about BUGS software in ecological modeling. We would again encourage you to become an active member of that community.