

Contents

<i>Preface</i>	ix
Why This Book?	ix
Goals	xi
Approach	xi

I

Fundamentals	1
1 PREVIEW	3
1.1 A Line of Inference for Ecology	4
1.2 An Example Hierarchical Model	11
1.3 What Lies Ahead?	15
2 DETERMINISTIC MODELS	17
2.1 Modeling Styles in Ecology	17
2.2 A Few Good Functions	21
3 PRINCIPLES OF PROBABILITY	29
3.1 Why Bother with First Principles?	29
3.2 Rules of Probability	31
3.3 Factoring Joint Probabilities	36
3.4 Probability Distributions	39
4 LIKELIHOOD	70
4.1 Likelihood Functions	70
4.2 Likelihood Profiles	73
4.3 Maximum Likelihood	75
4.4 The Use of Prior Information in Maximum Likelihood	76
5 SIMPLE BAYESIAN MODELS	78
5.1 Bayes' Theorem	79
5.2 The Relationship between Likelihood and Bayes'	84
5.3 Finding the Posterior Distribution in Closed Form	85
5.4 More about Prior Distributions	88

6	HIERARCHICAL BAYESIAN MODELS	105
6.1	What Is a Hierarchical Model?	106
6.2	Example Hierarchical Models	107
6.3	When Are Observation and Process Variance Identifiable?	139



Implementation 141

7	MARKOV CHAIN MONTE CARLO	143
7.1	Overview	143
7.2	How Does MCMC Work?	144
7.3	Specifics of the MCMC Algorithm	148
7.4	MCMC in Practice	175
8	INFERENCE FROM A SINGLE MODEL	178
8.1	Model Checking	178
8.2	Marginal Posterior Distributions	187
8.3	Derived Quantities	191
8.4	Predictions of Unobserved Quantities	193
8.5	Return to the Wildebeest	197
9	INFERENCE FROM MULTIPLE MODELS	205
9.1	Model Selection	206
9.2	Model Probabilities and Model Averaging	218
9.3	Which Method to Use?	223



Practice in Model Building 225

10	WRITING BAYESIAN MODELS	227
10.1	A General Approach	227
10.2	An Example of Model Building: Aboveground Net Primary Production in Sagebrush Steppe	231
11	PROBLEMS	237
11.1	Fisher's Ticks	238
11.2	Light Limitation of Trees	239
11.3	Landscape Occupancy of Swiss Breeding Birds	240
11.4	Allometry of Savanna Trees	241
11.5	Movement of Seals in the North Atlantic	242

12 SOLUTIONS	244
12.1 Fisher's Ticks	244
12.2 Light Limitation of Trees	249
12.3 Landscape Occupancy of Swiss Breeding Birds	252
12.4 Allometry of Savanna Trees	257
12.5 Movement of Seals in the North Atlantic	261

IV

Advanced Topics 265

13 Missing Data	267
13.1 Unintentionally Missing Data	268
13.2 Data Missing by Design	277
13.3 Guidance	288
14 Spatial Models	289
14.1 Introduction	289
14.2 Continuous and Discrete Spatial Models	289
14.3 Spatial Point Process Models	302

<i>Afterword</i>	309
Reviewing Bayesian Models in Papers and Proposals	309
Continuing Your Learning	311

<i>Acknowledgments</i>	315
------------------------	-----

<i>Appendix A: Probability Distributions and Conjugate Priors</i>	317
---	-----

<i>Bibliography</i>	323
---------------------	-----

<i>Index</i>	335
--------------	-----

1 Preview

Art is the lie that tells us the truth.

—*Pablo Picasso*

All models are wrong but some are useful.

—*George E. P. Box*

Pablo Picasso was a contemporary of George Box, a statistician who had an enormous impact on his field, writing influential papers well into his 90s. Both sought to express truth about nature, but they used tools that were dramatically different — Picasso brushing strokes on canvas, and Box writing equations on paper. Given the different ways they worked, it is remarkable that Box and Picasso made such similar statements about the central importance of abstraction to insight. Abstraction plays a role in all creative human enterprise — in art, music, literature, engineering, and science. We create abstractions because they allow us to focus on the most important elements of a problem, those relevant to the objectives of our work, without being distracted by elements that are not relevant.

Scientific models are, above all else, abstractions. They are statements about the operation of nature that purposefully omit many details and thus achieve insight that would otherwise be discursively obscured. They provide unambiguous statements of what we believe is important. A key principle in modeling and statistics — in science, for that matter — is the need to reduce the dimensions of a problem. A data set may contain a thousand observations. By reducing its dimensions to a model with a few parameters, we are able to gain understanding.

However, because models are abstractions and reduce the dimensions of a problem, we must deal with the elements we choose to omit. These elements create uncertainty in the predictions of models, so it follows that assessing uncertainty is fundamental to science. Scientists, journalists, logicians, and attorneys alike can rightly claim to make statements based on evidence, but only scientific statements include evidence tempered by uncertainty quantified. We know what is certain only to the extent that we can say, with

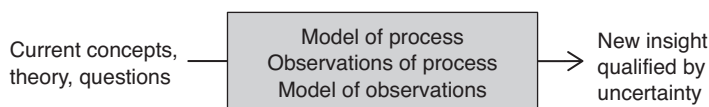


Figure 1.1.1. The fundamental challenge in ecological research is to establish a credible line of inference extending from concepts and theory to new insight tempered by uncertainty.

confidence, what is uncertain. Sharpening our thinking about uncertainty and learning how to estimate it properly is a main theme of this book.

Your science will have impact to the extent that you are able to ask important questions and provide compelling answers to them. Doing so depends on establishing a line of inference that extends from current thinking, theory, and questions to new insight qualified by uncertainty (fig. 1.1.1). This book offers a highly general, flexible approach to establishing this line of inference. We cannot help you pose novel, interesting questions, but we can teach an approach to inference applicable to an enormous range of research problems, an approach that can be understood from first principles and that can be unambiguously communicated to other scientists, managers, and policy makers. We emphasize that understanding the principles of this framework allows you to customize your analyses to accommodate the inevitable idiosyncrasies of specific problems in research.

We sketch that framework in this chapter to give a general sense of where this book is headed, a preview we use to motivate the development of concepts and principles in the chapters that follow. There should be details of our approach that are unfamiliar, otherwise you probably don't need this book. Soon enough, we will explain those details fully. For now, we offer a somewhat abstract overview followed by a concrete example as an enticement to read on. It will be rewarding to return to this section after you have worked through the book. We hope you will be pleasantly surprised by your increased understanding of our small preview. The only part of this chapter that is essential for the remainder of the book is understanding our notation, which we describe in section 1.1.1.

1.1 A Line of Inference for Ecology

Virtually all research problems in ecology share a set of features. We want to understand how the state of an ecological system changes over time, across space, or among individuals. We seek to understand why those changes occur. Our understanding usually depends on a sample drawn from all

possible instances of the state because we want to make statements about a system that is too large to observe fully. The observations in that sample are often related imperfectly to the true state. In the subsections that follow, we lay out an approach first described by Berliner (1996)¹ for modeling the imperfect observations that arise from a process we want to understand. It does not apply to all research problems, but it is sufficiently general and flexible that it applies to most.

1.1.1 Some Notation

Before we proceed, we must introduce some notation. Boldface lower-case letters will indicate vectors (e.g., θ , $\mathbf{a}$), and lightface lowercase letters, scalars (θ , a).² Bold capital letters will be used for matrices (e.g., $\mathbf{A}$). The symbol θ will indicate a vector of parameters, and, of course, θ will indicate a single parameter.³ The letter $\mathbf{y}$ will indicate a vector of data, $\mathbf{Y}$ a matrix, and y or y_i a single observation. Corresponding notation using $\mathbf{x}$, x , and x_i will be used for predictor variables, also called covariates. The notation $[a|b, c]$ will be used for the probability distribution of the random variable a conditional on the parameters b and c .⁴ Deterministic models will be denoted by $g()$ with arguments necessary for the model within the parentheses. Notation will be added as needed, in context.

1.1.2 Process Models

Process models include a mathematical statement that depicts a process and a way to account for uncertainty about the process. To compose a process model we start by thinking about the true state (z) of an ecological system. That state could be the size of a population, the flux of nitrogen from the soils of a grassland, the number of invasive plants in a community, or the area of landscape annually disturbed by fire. We seek to understand influences on that state, the things that cause it to change. We write an equation,⁵ a

¹For elaboration, see Wikle (2003); Clark (2003b); Cressie et al. (2009), and Wikle et al. (2013).

²It might be useful to review scalars, vectors, and matrices. Most ecologists are familiar with matrices — rows and columns of numbers. Vectors are “one dimension” of a matrix, that is, a row or a column. Scalars are a single element. So, a matrix might be $\mathbf{A} = \begin{pmatrix} c & d \\ d & f \end{pmatrix}$; a vector, $\mathbf{a} = (c, d)'$; and a scalar, c . We will use the notation $()'$ to list the elements of a vector.

³We illustrate with the familiar example $y = \beta_0 + \beta_1 x$. The vector of parameters from the model is $\boldsymbol{\beta} = (\beta_0, \beta_1)'$. When discussing parameters generically, we will use θ .

⁴The symbol “[|]” reads conditional on.

⁵Or equations. To keep things simple, we focus on a single equation here, but we might build a system of equations making multiple predictions of states.

deterministic model that represents our ideas about the behavior of the state of interest and the quantities that influence it.⁶ When we say the model is deterministic, we mean that for a given set of parameters and inputs, it will make precisely the same predictions. We use the notation $g(\theta_p, \mathbf{x})$ to represent the deterministic part of a process model, where $g()$ is any mathematical function, θ_p is a vector of parameters in the model, and $\mathbf{x}$ is one or more explanatory variables that we hypothesize influence the true state.

Our deterministic model is an abstraction, so it follows that we have omitted influences on the true state from the model, and we must deal with our omissions. If we model aboveground net primary production of a grassland as a function of growing season precipitation, we have brushed aside the influence of grazing intensity and precipitation that occurs during the dormant season. If we model reproductive success of individuals as a function of age and genotype, we have ignored variation contributed by their nutritional status. A model of harvest from a fishery based on observations of stock size and sea temperature omits the effect of variation in the food web. We recognize that these neglected influences shape the behavior of the true state by treating them stochastically, by estimating a parameter, σ_p^2 , that subsumes all the unmodeled influences on the true state. Including this stochastic component allows us to estimate a statistical distribution (fig. 1.1.2A) for the true state,

$$\underbrace{[z|g(\theta_p, \mathbf{x}), \sigma_p^2]}_{\text{process model}}, \quad (1.1.1)$$

where the bracket notation $[z|g(\theta_p, \mathbf{x}), \sigma_p^2]$ means the distribution of z conditional on $g(\theta_p, \mathbf{x})$ and σ_p^2 .⁷ If the notation is somewhat unfamiliar at this point, don't worry; equation 1.1.1 simply says that if we know the functional form $g()$ and the values of θ_p , $\mathbf{x}$, and σ_p^2 , we can specify the probability distribution of the true state, z (fig. 1.1.2 A).

We want to determine the probability distribution of the true state as well as the probability distributions of the parameters in our model. Doing

⁶As a tangible example, we might model the influence of phytoplankton biomass (x) on zooplankton biomass (z) with a linear model, $z = \theta_0 + \theta_1 x$. In this case $z = g(\theta, x) = \theta_0 + \theta_1 x$.

⁷This notation was first introduced by Gelfand and Smith (1990) as a way to reduce clutter. It has become widely used by statisticians and ecologists. There is a small caveat needed here. We are using the arguments for probability distributions as if they were the mean ($g(\theta_p, x)$) and variance (σ_p^2) of the distributions. The arguments for most of the statistical distributions are parameters that are functions of the mean and variance, a topic we treat in detail in section 3.4.4. For now, we take the liberty of treating the mean and variance as the needed parameters because they are familiar to ecologists.

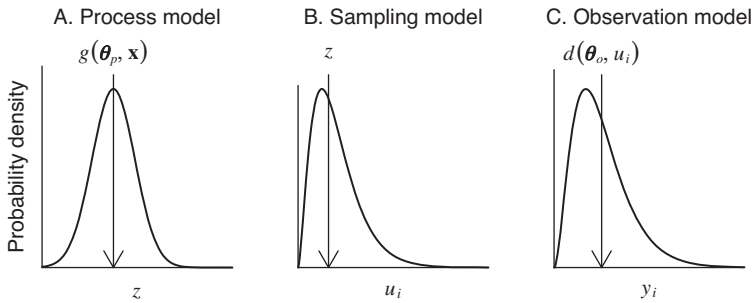


Figure 1.1.2. An observation on the operation of an ecological process (y_i) is linked to ideas about how the process works via three linked probability distributions. Hierarchical models consist of linked distributions. As we move from panel **A** to **C**, notice how the random variables (on the x -axis) becomes the mean (or other central tendency) of the distribution in the next panel. **(A)** The deterministic model, $g(\theta_p, \mathbf{x})$ predicts a true state, z , as a function of parameters θ_p and predictor variables $\mathbf{x}$. There is uncertainty in the predictions, σ_p^2 , that arises because there are influences on z that are not represented in the model. Thus, the distribution in **A** is $[z|g(\theta_p, \mathbf{x}), \sigma_p^2]$. If the deterministic model predicts z well, then the distribution in **A** shrinks toward the $g(\theta_p, \mathbf{x})$ arrow. **(B)** There are almost always more instances of z in nature than we can hope to observe. Observations of individual instances of z define a sampling distribution, $[u_i|z, \sigma_s^2]$, the breadth of which depends on σ_s^2 . As variation among observations declines, the distribution in **B** shrinks toward the z arrow. Note that the mean of the distribution shown by the arrow is not at the peak of the distribution, because the distribution is skewed. **(C)** Observations (y_i) of instances of the true state u_i are often biased, such that $y_i \neq u_i$. A deterministic observation model, $d(\theta_o, u_i)$, corrects for this bias. Uncertainty in this correction (σ_o^2) leads to the observation distribution $[y_i|d(\theta_o, u_i), \sigma_o^2]$. As uncertainty in the observation model declines, the distribution in **C** shrinks toward the $d(\theta_o, u_i)$ arrow.

so requires evaluating the predictions of the process model against data. The data can be obtained in experiments or observational studies; they can be measurements we plan to collect or have already collected. This linkage between process models and observations is discussed next.

1.1.3 Sampling Models

We can rarely observe all instances of the true state in the system we study. Instead, we take a sample of $i = 1, \dots, n$ observations of the true state, and we notate the i th observation as u_i . This sample might be biomass from plots on a grassland landscape where we seek to understand the true state, above-ground productivity. It might be presence or absence of an exotic fungus on

trees in a stand where we want to understand infestation of the stand. It might be classifications of zooplankton in aliquots from a stream where we want to estimate the stream's species richness. Uncertainty arises because our sample assuredly will not represent the true state perfectly. Again, we represent this uncertainty stochastically using a probability distribution relating the true state to an observation

$$\underbrace{\left[u_i | z, \sigma_s^2 \right]}_{\text{sampling model}}, \quad (1.1.2)$$

where σ_s^2 represents sampling variation (fig. 1.1.2 B). Expression 1.1.2 implicitly assumes that we can observe instances of the true state without bias, which simply says that if we collect many observations, then the average (i.e., also called the *expected value* of the observations, $E(u)$) of the observations equals the mean of the distribution of the true state, $E(u) = z$. (Expectation will be treated in detail in chapter 3.) We realize samples of the true state is a nuanced concept — soldier on, things will become clear in the next section.

1.1.4 Observation Models

The assumption that we can observe the true state perfectly may not be reasonable. When we count animals, some are overlooked. When we use Lidar⁸ to estimate the heights of 10,000 trees, we do not measure the height of each tree using a ladder and a meter tape (thankfully) but instead observe backscatter from a laser beam. When we measure nitrogen mineralization, we do not follow the fate of individual nitrogen atoms but measure the net change in the extractable soil ammonium pool over time. The mismatch between what we observe and the true state requires a model of the observations, which we notate as $d(\theta_o, u_i)$, where θ_o are parameters. It is important to understand that u_i is the quantity we would observe if we could *perfectly* observe the instance of the true state in a draw from all the instances, without any bias injected by our observation process. We use y_i to notate the actual measurements we have in hand, including error resulting from the way we observe the u_i . The observation model serves to eliminate the bias found in our observations, y_i relative to an instance (u_i) of the true state drawn from the distribution of z . The probability distribution of the observations (fig. 1.1.2 C) arising from the observed instances of the

⁸Lidar, light detection and ranging, a technique used in remote sensing.

true state is

$$\underbrace{\left[y_i | d(\theta_o, u_i), \sigma_o^2 \right]}_{\text{observation model}}, \quad (1.1.3)$$

where the σ_o^2 represents all the influences on the y_i that are not represented in $d(\theta_o, u_i)$. As a simple example, the σ_o^2 could be the variance of the predictions of a regression model used to calibrate observations against true values. We emphasize that model $d(\theta_o, u_i)$ is needed to offset bias in the y_i . If our observations are unbiased, then there is no need for an observation model (i.e., $y_i = u_i$), and the uncertainty in our data arises solely from sampling variation (i.e., eq. 1.1.3). For simplicity, we have ignored the sampling variation and observation errors that might influence the $\mathbf{x}$, but these could be handled in the same way as we have done for the $\mathbf{y}$ (i.e., we would use probability distributions like eqs. 1.1.2 and 1.1.3).

1.1.5 Parameter Models

Because the approach we sketch is Bayesian, we also require models of the parameters expressing what we knew about the parameters when we began our investigation, that is, our *prior* knowledge. This knowledge is expressed in probability distributions, one for each parameter we seek to estimate⁹

$$\underbrace{\left[\theta_p \right] \left[\theta_o \right] \left[\sigma_p^2 \right] \left[\sigma_s^2 \right] \left[\sigma_o^2 \right]}_{\text{parameter models}}. \quad (1.1.4)$$

These distributions must have numeric arguments that specify our current knowledge of the probability distribution of the parameters. The arguments can be chosen to make the distributions informative or vague, but as you will see, we will encourage you to make priors as informative as knowledge and scholarship allow. We might know a lot about a parameter or we might know very little.

1.1.6 The Full Model

We are now equipped to write a mathematical expression representing our ideas about the operation of an ecological process linked to data in a way

⁹Equation 1.1.4 requires the assumption that the parameters $\theta_p, \theta_o, \sigma_p^2, \sigma_s^2$, and σ_o^2 are statistically independent. We will explain this idea in greater detail later.

that includes all sources of uncertainty — in the process, in our sample of the process, and in the way we observe it (fig. 1.1.2). Given a single observation y_i , we write a posterior distribution as

$$\left[\underbrace{z, \theta_p, \theta_o, \sigma_p^2, \sigma_s^2, \sigma_o^2, u_i}_{\text{unobserved}} \mid \underbrace{y_i}_{\text{observed}} \right] \propto \underbrace{y_i | d(\theta_o, u_i), \sigma_o^2}_{\text{observation model}} \underbrace{u_i | z, \sigma_s^2}_{\text{sampling model}} \underbrace{z | g(\theta_p, \mathbf{x}), \sigma_p^2}_{\text{process model}} \underbrace{[\theta_p] [\theta_o] [\sigma_p^2] [\sigma_s^2] [\sigma_o^2]}_{\text{parameter models}}. \quad (1.1.5)$$

Equation 1.1.5 is Bayesian and hierarchical.¹⁰ It is *Bayesian* because it treats the unobserved quantities as random variables. This treatment allows us to make statements about the probability distributions of all of the unobserved quantities based on the observed ones. It is *hierarchical* because the z and u_i are found on both sides of a conditioning symbol “ $\mid$ ”, illustrating a powerful tool for simplifying problems that we will discuss more fully soon. The observation model includes our knowledge of the relationship between the true state and our observations of it and the uncertainty that occurs because that relationship is imperfect. The sampling model includes the uncertainty that comes from observing a subset of instances of the true state. The process model represents our hypotheses about the ecological process by specifying a probability distribution defining our knowledge and our uncertainty about the true state and the factors that control its behavior. The parameter models allow us to exploit previous estimates of parameters that we have made ourselves or that have been made by others. Together these models provide a line of inference extending from concepts to insight for a broad range of research problems (fig. 1.1.1). We can use equation 1.1.5 to obtain estimates of unobserved states, parameters, and quantities of interest derived from parameters and states. All of these estimates are properly tempered by uncertainty in a statistically coherent way.

In the remainder of this book, we develop the principles needed to understand equation 1.1.5 and to apply it to research problems in ecology. We tailor it to match the needs of the particular problem at hand. But first, we provide an example of its use.

¹⁰The symbol $\propto$ means “is proportional to.” It may be confusing at this point that the observed $\mathbf{x}$ does not show up in the list of observed quantities on the left-hand side of the proportionality. Bear with us. We will make that clear as we proceed.

1.2 An Example Hierarchical Model

We are now going to apply the general framework we have previewed to a specific problem. Remember that we don't expect this section to be familiar to you. Although the application here focuses on learning about population dynamics of a large mammal from a time series of observations in East Africa, it could be about any topic, any research design, and any location. We urge you to draw analogies to your own work as the example develops.

The Serengeti wildebeest (*Connochaetes taurinus*) population migrates across the grasslands of Tanzania and Kenya in an annual cycle driven by availability of green plant biomass (Boone et al., 2006). During the late 1800s, the viral disease rinderpest created a panzootic, decimating the wildebeest and other wild and domestic ruminants. Numbers of these animals remained low until the 1950s when a campaign to vaccinate cattle in the pastoral lands surrounding the Serengeti eliminated the virus in wildlife (Plowright, 1982). Survivorship of wildebeest, particularly juveniles, doubled after rinderpest was eradicated, and the population grew rapidly until the mid 1970s when density-dependent mortality produced a quasi-equilibrium. This steady state appears to have been caused by intraspecific competition for green plant biomass during the dry season. This biomass varies in approximate proportion to annually variable rainfall (Mduma et al., 1999).

Developing informative models depends on a clearly stated question. An unambiguous question is vital because it guides the abstraction that we formulate; it informs our decisions about which variables and parameters we need to include in our model. In this particular example we ask the question: How does variation in weather modify feedbacks between population density and population growth rate in a population of large herbivores occupying a landscape where precipitation is variable in time? Answering this question requires a model that portrays density dependence, effects of precipitation, and their interaction.

1.2.1 Process Model

We pose the deterministic model portraying growth of the wildebeest population as

$$N_t = g(\beta, N_{t-1}, x_t, \Delta t) = N_{t-1} e^{(\beta_0 + \beta_1 N_{t-1} + \beta_2 x_t + \beta_3 N_{t-1} x_t) \Delta t}, \quad (1.2.1)$$

where N_t is the unobserved, true abundance of wildebeest in year t (corresponding to z in eq. 1.1.5); x_t is a measure of the annual rainfall

that influences population growth during $t - 1$ to t ; and $\Delta t = 1$ year.¹¹ The meaning of the coefficients β requires some thought. A critical part of gaining insight from models and data is careful thinking about the biological meaning of the parameters in our models and their relationship to existing concepts and theory, the need for which is illustrated here.

We first consider the intercept, β_0 , a logical place to start. If the units of precipitation, x_t , are centimeters per year, ranging from 0 to a large number, then algebra dictates that $e^{\beta_0 \Delta t}$ is the proportional change in the population during the period t to $t + 1$ that occurs when abundance is zero and rainfall is zero. The abundance equals zero part might seem a bit odd at first glance, because there is no population to grow when $N_t = 0$, but the parameter nonetheless serves to define the upper limit on population growth rate as numbers approach zero. There is no problem here. But the part about zero rainfall is troubling because it makes β_0 difficult to interpret biologically or at least makes it somewhat unuseful in its biological interpretation. Why would we want to estimate a parameter determining population growth when the population is low and rainfall is low? Alternatively, if we define x_t in terms of divergence from the long-term average by subtracting the mean rainfall from each year's observation, then the definition of β_0 becomes more sensible. In this case, $e^{\beta_0 \Delta t}$ is the proportional change in population size that occurs when the population is zero and rainfall is average. This definition allows us to relate our model to well-established theory; β_0 is analogous to the intrinsic rate of increase in the logistic equation.¹²

Once we have defined the intercept in a sensible way, the slope terms are easily interpreted. The parameter β_1 represents the strength of density dependence, that is, the change in the per capita population growth rate that occurs for each animal added to the population.¹³ Again relating our model to classical theory, $\beta_1 = -(r_{\max}/K)$, where K is the carrying capacity, that is, the population size at which the long-term average population

¹¹The identical form $\log(N_t) = \log(N_{t-1}) + (\beta_0 + \beta_1 N_{t-1} + \beta_2 x_t + \beta_3 N_{t-1} x_t) \Delta t$ might be more familiar to ecologists accustomed to a general linear modeling framework.

¹²This might require a bit of explanation. Equation 1.2.1 can be rearranged to $\log(N_t/N_{t-1}) = (\beta_0 + \beta_1 N_{t-1} + \beta_2 x_t + \beta_3 N_{t-1} x_t) \Delta t$. If we drop the $\beta_2 x_t$ and $\beta_3 N_{t-1} x_t$ terms and assume Δt is small, the resulting expression approximates $1/N \cdot dN/dt = r - rN/K$, where N is the population size, r is the intrinsic rate of increase, and K is the carrying capacity — the long-term average population size at which $1/N \cdot dN/dt = 0$. Thus, the portion of our model representing density dependence is a version of the logistic equation in discrete time, also known as the *Ricker equation*, where $\beta_0 = r$, and $\beta_1 = r/K$.

¹³You might be tempted to make the sign for β_1 negative, because density dependence reduces the per capita population growth rate as population size increases. There is a good reason to use addition. When the model is fit to data, we want to estimate the sign and the value of β_1 without needing to “back transform” the sign. Always use addition in additive models that you will fit to data.

growth rate equals zero. The parameter β_2 expresses the strength of the effect of variation in rainfall, that is, the change in the per capita population growth rate per unit deviation from the long-term mean. Finally, the parameter β_3 determines the magnitude of the effect of rainfall on the effect of density.

Estimating the β s will inform the question we posed, but we don't pretend that they capture all the influences on wildebeest population dynamics. We recognize that there are other processes, for example, predation, poaching, and disease, that shape wildebeest numbers over time. We can acknowledge these other processes exist without portraying them explicitly. Instead, we lump these unmodeled effects into a single parameter, σ_p^2 . The model of the process for a population at time t including deterministic and stochastic components is

$$\left[N_t | g(\beta, N_{t-1}, x_t, \Delta t), \sigma_p^2 \right]. \quad (1.2.2)$$

1.2.2 Sampling Model

The wildebeest population was estimated¹⁴ on 20 occasions during 1961–2008 using spatially replicated counts of animals on georectified aerial photographs arrayed along transects, with each photograph covering a known area (Norton-Griffiths, 1973). For simplicity, we ignore the aspect of sampling contributed by the transects and treat the photographs as if they were a random sample from the area used by the population. Thus, we assume that each photograph provided a statistically independent estimate of population density calculated as the observed count divided by the area covered by the photograph. The stochastic sampling model representing the relationship between these observations and the true population size is

$$\left[y_{tj} \left| \frac{N_t}{a}, \sigma_s^2 \right. \right], \quad (1.2.3)$$

where y_{tj} is the density of animals at year t on photograph j , and a is the total area from which the sample of photographs was drawn — the area used by the population we are modeling — which we assume to be a known constant. Note that equation 1.2.3 implies that there is no bias in the estimate of animal density on a given photograph, which is the assumption made by the researchers who collected the data. Although this assumption is

¹⁴Portions of this time series have been published in Hilborn and Mangel (1997) and Mduma et al. (1999). The most recent data were graciously provided by A.R.E. Sinclair, Ray Hilborn, and Grant Hopcraft. We use these data later to illustrate making inference from a single model.

reasonable for large animals counted in open habitat, we could obviate the need for this assumption by modeling the animals that were present and not observed.

As a purely hypothetical example, imagine that we fitted a sample of animals with high-resolution telemetry instruments that would allow us to know whether each animal was within a photograph. Imagine also that we marked them with highly visible neckbands to allow us to distinguish between animals that were fitted with instruments and those that were not. We could use these observations to estimate the probability (ψ) that an animal truly present was counted. In this case our combined observation and sampling model would be

$$\underbrace{\left[y_{ij} | \psi n_{ij}, \sigma_o^2 \right]}_{\text{observation model}} \underbrace{\left[n_{ij} \left| \frac{N_t}{a}, \sigma_s^2 \right. \right]}_{\text{sampling model}}, \quad (1.2.4)$$

where n_{ij} are animals that are truly present on photograph j during year t , and σ_o^2 is the uncertainty associated with the estimate of ψ . We could also deal explicitly with the likely cases in which it was not certain whether an animal was on a sampled photograph.

To simplify the example and to exploit published data, we make the heroic assumption that rainfall at time t , x_t , is measured without error; we are treating it as known, but we are not obliged to do so. We could develop sampling models and observation models for the rainfall data in the same way we did for the count data (i.e., eqs. 1.2.3 and 1.2.4).

1.2.3 Full Model

Combining the sampling model (eq. 1.2.3) and process model (eq. 1.2.2) with models for the parameters, we obtain the foundation for a full Bayesian analysis capable of estimating the parameters and the unobserved, true population size (fig. 1.2.1); however, work remains to be done. The full model written thus far (fig. 1.2.1) applies to a single year of observations and a single photograph, but, of course, we want to use all the years and all the observations within a year. We have not yet chosen specific probability distributions for the stochastic components, and we must do so in a sensible way. We may need to deal with correlation among photographs or among annual estimates in the time series of data, and with errors in the estimation of rainfall. We need to lay out a method for numerically estimating the parameters and unobserved states. However, although tasks remain, all proceed in a logical way from the foundation we have built (fig. 1.2.1). All

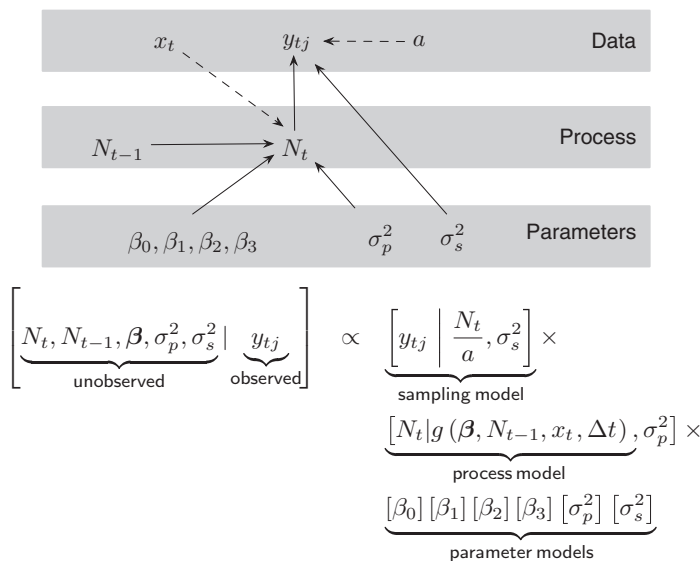


Figure 1.2.1. Hierarchical Bayesian model of the dynamics of the Serengeti wildebeest population. The true population size at time t is modeled using the deterministic model $g(\beta, N_{t-1}, x_t, \Delta t) = N_{t-1} e^{(\beta_0 + \beta_1 N_{t-1} + \beta_2 x_t + \beta_3 N_{t-1} x_t) \Delta t}$, which represents the effects of the true, unobserved population size (N_t); the dry season rainfall (x_t); and their interaction on population growth. The biological interpretation of the parameters β are given in the text. The parameter σ_p^2 represents all the effects on wildebeest abundance not represented by the β . The y_{tj} is a single observation of population density obtained from one of $j = 1, \dots, n_t$ photographs chosen from area $= a \text{ km}^2$. The parameter σ_s^2 represents sampling error. The arrow diagram is a Bayesian network, also called a *directed acyclic graph*, that can be used as a visual guide for properly writing the full model. The solid lines show stochastic relationships. The dashed lines show deterministic relationships, implying that the quantities at the tails of the arrows are known without error. Instructions for drawing diagrams like this one will be developed in subsequent chapters.

are manageable. Completing them allows us to reliably estimate unobserved states, parameters, and derived quantities of interest.

1.3 What Lies Ahead?

One of the aims of this book is to enable ecologists to write equations for models that allow data to speak informatively. Our goal is to provide an understanding of principles needed to accomplish this vital task. We will show how an enormous range of research problems in ecology can be

decomposed into a set of sensible parts, as we have illustrated. The specific example we offered no doubt raised more questions than it answered, which was our intention. Notable among these questions might be the following:

1. Why does the Bayesian approach to analysis work? What are the probabilistic foundations of inference from models like this one?
2. How does a Bayesian analysis relate to analyses based on maximum likelihood?
3. How do we choose the appropriate statistical distributions for the stochastic components of the model (fig. 1.2.1)?
4. How can we incorporate multiple sources of data in estimates of parameters and states?
5. How do numerical methods allow us to implement the model and obtain results? How do these methods work?
6. How can we evaluate the model to assure it adequately represents the data?
7. What can we conclude about the operation of the ecological process based on estimates of the parameters and states? What do we do about derived quantities; for example, how do we make inferences about the population size where growth rate is maximum?
8. What if we want to evaluate the strength of evidence for this model (eq. 1.2.1) relative to competing ones, for example, a model with nonlinear density dependence?
9. How do we predict future states with honest estimates of uncertainty?

Answering these questions in a clear and accessible way for a broad range of research problems in ecology is the goal of remainder of this book.

INDEX

- accept-reject sampling, 148; efficiency of, 155
- AIC, *See* Akaike's information criterion
- Akaike's information criterion, 210–212
- Bayes factors, 220–221
- Bayes' theorem, xii, 32, 35, 78–86, 105, 106, 219, 221; derivation, 79–82; joint distribution, 82; marginal distribution of data, 82; normalizing the joint distribution, 83; posterior distribution, 83; prior distribution, 82
- Bayesian inference: difference from other branches of statistics, 78; from a single model, 78–205; from multiple models, 205–225; known vs unknown, 79; observed vs unobserved, 79; relationship to likelihood, xii, 84–85, 104; unobserved quantities as random variables, 78
- Bayesian information criterion, 211
- Bayesian model averaging, 218–231
- Bayesian network, xiv, 15, 36–39, 105, 109, 117, 121, 126, 131, 132, 175, 210, 229–231, 234, 245, 250, 251, 253, 255, 258, 259; basis for factoring joint distribution, 38, 107, 110, 111, 114, 231, 255; common mistake when drawing, 255; for full-conditional distributions, 158
- Bayesian p-value, 185–187, 200, 246
- Bernoulli distribution, 55, 61, 69, 89, 90, 94, 95, 217, 252, 254, 255, 257
- beta distribution, 61, 67, 71, 86, 87, 89, 91, 100, 102, 130, 247, 251; moment matching, 67, 71, 239, 251
- BIC, *See* Bayesian information criterion
- binomial distribution, 54, 87, 90, 100, 128, 129, 248, 257
- BMA, *See* Bayesian model averaging
- borrowing strength, 124–125
- Box, George E. P., 3
- bracket notation, 5, 6
- burn-in, 168
- CDF, *See* cumulative distribution function
- central limit theorem, 56, 58, 256
- collaboration with statisticians, x, 4, 175, 312
- composition sampling, 181, 194, 196
- conditional probability, 32, 34, 39, 73, 80, 180; definition, 34; notation, 32, 72
- conjugate prior distribution, 85–89, 93, 126, 128, 164, 317; beta binomial, 86, 128; gamma Poisson, 244; lognormal variance, 126, 235, 248, 258; normal mean, 166; normal variance, 95, 167, 262
- continuous random variable, definition, 40
- convergence of MCMC chain, 167–172; burn-in, 168; Gelman-Rubin diagnostic, 170–172; Geweke diagnostic, 172; Heidelberg-Welch test, 172; Raftery-Lewis diagnostic, 172; visual inspection of trace plots, 170
- covariance, 63
- covariance matrix, 51, 63, 104, 136, 137, 261
- covariate, 5, 22, 79, 114, 122, 125, 129, 131, 135, 186, 217, 229, 230, 239, 241, 257; standardized, 256; uncertainty in, 118, 133, 229, 249–251
- credible interval: equal-tailed, 189; highest posterior density, 188–191
- cross-validation, 207–208
- cumulative distribution function, 46–48, 201
- data model, 9, 72, 93, 100, 129, 131, 133–135, 198, 240, 243
- data simulation for verifying model coding, 173

336 • Index

- derived quantity, 191–192, 201, 236, 251, 256, 260
- detection probability, 62, 128, 240, 253, 256
- deterministic model, 7, 15–239;
 - definition, 21; exponential, 11, 15, 23, 116, 248; Gompertz, 27; Hill, 27;
 - inverse logit, 23, 64, 66, 254, 255;
 - linear, 232, 234; Michaelis–Menten, 26, 161, 175, 249; monomolecular, 27;
 - power function, 26, 131, 257, 258;
 - simple linear, 22; threshold, 28
- deviance, 212
- deviance information criterion, 212
- DIC, *See* deviance information criterion
- dimension reduction, 29
- Directed acyclic graph, *See* Bayesian network
- Dirichlet distribution, 63
- discrete random variable, definition, 40
- distribution, 38–63; arguments, 50;
 - Bernoulli, 55, 61, 69, 89, 90, 94, 95, 217, 252, 254, 255, 257; beta, 61, 67, 71, 86, 89, 91, 100, 102, 130, 247, 251;
 - binomial, 54, 87, 90, 100, 128, 129, 248, 257; cumulative distribution function, 47–48; Dirichlet, 63;
 - expected value, 44; exponential, 59;
 - gamma, 42, 56–59, 65, 66, 68, 88, 111–112, 115; inverse gamma, 61, 95, 112, 115, 126, 235, 248; lognormal, 56–58, 115, 125, 126, 162, 233, 248, 257; marginal, 48–50; mixture, 67–69, 134; moment matching, 64–67;
 - moments, 44, 46; multinomial, 55;
 - multivariate lognormal, 137;
 - multivariate normal, 62, 135, 261;
 - negative binomial, 54, 59, 69, 111, 118, 246; normal, 56, 115, 132, 198, 249; notation, 41–42; Poisson, 52–139;
 - probability density function, 42–43;
 - probability mass function, 40; quantile function, 47; uniform, 62
- distribution function, *See* cumulative distribution function
- effective number of parameters, 212, 214
- equivariance property of MCMC, 191
- error in variables model, 242
- example models: above ground net
 - primary production, 231–236;
 - allometry of savanna trees, 241–242, 257–260; arctic copepods, 87; effect of functional traits on tree growth, 131–135; fecundity of spotted owls, 109–122; landscape occupancy of breeding birds, 240–241, 251–257; light limitation of tree growth, 249–251; moth predation, 100–102; mussel diversity, 69; Nitrous oxide emissions from agricultural soils, 122; seal movement, 242; Serengeti wildebeest, 197–204; sexual dimorphism, 68; stage structured population, 136–139, 238, 239, 244–249; ticks on sheep, 238, 239, 244–249; treesnake predation on lizards, 127–131; whitebark pine infection, 86
- expected value, 8, 44; definition, 45
- exponential distribution, 59; relationship to compartment models, 60; relationship to individual based models, 60
- exponential model, 11, 15, 23, 116, 248
- factoring joint distribution, 1, 36–39, 105–111, 114, 120–122, 139, 210, 217, 228, 231, 241, 263, 309
- fake data simulation, *See* data simulation for verifying model coding
- forecast, *See* predictive process distribution
- full-conditional distribution, 2, 125, 157–160, 163, 167, 310; basis for constructing MCMC algorithm, 160–163
- gamma distribution, 42, 56–59, 65, 66, 68, 88, 111–112, 115, 244; alternative parameterization, 59; moment matching, 65, 88, 112, 235, 247

- gamma function, 54
- Gelman and Rubin diagnostic, 200
- Gelman-Rubin diagnostic, 170–172
- Geweke diagnostic, 172
- Gibbs update, 163–167; normal mean and variance example, 165–167
- goals of book, xi
- Gompertz function, 27
- group-level effect, *See* group-level model; random effect
- group-level model, 113, 125, 127, 229, 234
- Heidelberg-Welch test, 172, 200
- hierarchical Bayesian model, 7, 105–140;
 - above ground net primary production example, 233–236; analysis of designed experiment, 127; borrowing strength, 124; breeding bird occupancy example, 240–241; common applications, 107; data model, 9, 129, 131–135, 198, 252, 261–263; definition, 105; factoring joint distribution into models of data, process, parameters, 108, 139; fecundity of spotted owls example, 122; general approach for writing, 227–231; group-level model, 122–124, 131, 234; hidden process, 127–131; landscape occupancy example, 251–257; latent state, 127, 138; light limitation of trees example, 239, 249–251; multiple sources of data, 135–139; Nitrous oxide example, 122–127; observation model, 229, 232–233; process model, 229, 234; sampling model, 229, 233; seal movement example, 242; separating observation from process variance, 139–140; Serengeti wildebeest, 10; Serengeti wildebeest example, 197–204; stage structured population example, 135–139; ticks on sheep example, 238, 239, 244–249; tree allometry example, 257–260; treesnake predation on lizards, 127–131; Tropical tree growth example, 131–135; uncertainty in covariate, 229
- highest posterior density interval, 3, 184, 189–191, 311
- Hill function, 27
- HPDI, *See* highest posterior density interval
- identifiability, 98, 104, 135, 139, 140, 170; of process and observation variance, 139–140
- indicator variable, 130, 234, 246
- indicator variable selection, 216–218
- individual variation, 30, 111–112, 246
- INLA, 2, 310
- integrating out, 50, 180, 185
- interpreting regression coefficients, 12
- inverse gamma, 61, 95, 112, 115, 126, 235, 248, 262
- inverse logit model, 23, 64, 66, 254, 255
- JAGS, x, 187, 199
- Jeffreys prior, 93
- joint distribution: as component of Bayes’ theorem, 83; factoring, 1, 36–39, 105–111, 114, 120–122, 139, 210, 217, 228, 231, 256, 263, 309
- joint prior distributions, 34, 217
- joint probability, 34
- latent state, 127, 131, 138, 157
- law of total probability, 35, 49
- likelihood, 77; function, 70–72; in Bayesian models, 72; maximum, 28, 76; principle, 75; prior information, 76–77, 209; profile, 73, 84; proportionality to probability, 72; ratio, 75; relationship to Bayesian inference, 84–85, 104; relativity of evidence, 75; use of term “support”, 75
- likelihood function, 70–72;
 - decomposition example, 71; definition, 72; tadpole example, 71
- likelihood profile, 73
- line of inference, 4–5

- linear model, 22
- link function, logit, 24
- local and proper scoring function, 213
- log predictive density, 207
- lognormal distribution, 56–58, 115, 125, 126, 162, 233, 248, 257
- marginal distribution of the data: as component of Bayes’ theorem, 83
- marginal distribution, explained, 48–50
- marginal probability of the data, *See* marginal distribution of data
- marginalization property of MCMC, 187
- Markov chain Monte Carlo algorithm, 28; accept-reject sampling, 148, 149; accumulation of random draws from posterior distribution, 141–177; autocorrelation of chain, 168; burn-in, 168; comparing posterior distribution with prior distribution, 201; composing, 156–157; composition sampling, 181, 194, 196; constructing from full-conditional distributions, 160–163; credible interval, 189, 191, 201; derived quantity, 191–193, 201, 260; equal-tailed credible interval, 189; equivariance property, 191; evaluating convergence, 2, 167–172, 199, 310; failure to converge, 170; full-conditional distribution, 157–160, 163; Gibbs sampling, 148; Gibbs updates, 148, 163–167; highest posterior density interval, 189–191; How does it work?, 144–148; hybrid updates, 167; implementing using software, 174–177; inference from output, 177–204; latent state, 201; marginal posterior distribution, 187–188, 199–204; marginalization property, 188; Metropolis ratio, 152–155; Metropolis updates, 151–155; Metropolis-Hastings updates, 155, 167; missing data, 199; mixing of chain, 170; moments of marginal posterior distribution, 184; output from, 183; posterior predictive distribution, 184, 185; predictions of unobserved quantity, 193–197; predictive process distribution, 194–197; proposal distribution, 149–151; sampling multiple parameters, 156–167; steps in, 147–148; summary statistics, 184, 188–189, 191; thinning chain, 170; verifying using simulated data, 172–174; What does it do?, 144–148
- Markov chain Monte Carlo: coding a sampler, 177
- maximum likelihood, 28, 76
- MCMC, *See* Markov chain Monte Carlo algorithm
- mean square prediction error, 206
- method of moments, *See* moment matching
- Metropolis updates, 151–155; ratio determining acceptance, 152–155
- Metropolis-Hastings updates, 155, 167; ratio determining acceptance, 155, 163
- Michaelis–Menten function, 26
- missing data, 199, 267–288; finite population inference, 286; ignorability, 278–279; cluster sample, 283–284; completely random sample, 281–283; randomized complete block example, 281; stratified random sample, 284–286; ignorable conditional on covariates, 280; in designed experiments and samples, 277–286; joint likelihood of missing and observed data, 269; missing at random, 270; missing completely at random, 269; missing covariates, 275–276; missing data imputation, 273–275; missing data model, 268; missing not at random, 271, 273–275; posterior predictive checks, 276–277
- missing data; finite population inference, 287
- missing data; ignorability; randomized complete block example, 280
- mixing MCMC, 168

- mixture distribution, 67–69, 224;
 - detection probability example, 128, 247, 252; female body mass example, 68; mussel example, 69; topical trees data model, 131; tropical trees data model, 134; zero-inflation, 128, 247
- mixture model, *See* mixture distribution
- model building process, xi, xiii
- model checking, 2, 178–187, 199, 223, 235, 245, 310; Bayesian p-value, 185–187, 200; posterior predictive distribution, 179–185
- model diagram, *See* Bayesian network
- model selection, 3, *See* multimodel inference, 256, 311; out-of-sample validation, 3, 311
- model selection uncertainty, 31
- model weights, 221
- modeling styles in ecology, 17, 18;
 - empirical, 19; simulation, 20; theoretical, 18
- models as abstractions, 3, 6
- models as hypotheses, 17
- moment matching, 64–67, 136, 228; beta distribution example, 66, 71, 247, 251; gamma distribution example, 65, 87; single parameter example, 67
- moments of distributions, 44–46
- monomolecular function, 27
- Monte Carlo integration, 46, 188, 192, 194, 197, 214
- MSPE, *See* mean square prediction error
- multilevel model, 122
- multimodel inference, 177–204; Akaike’s
 - information criterion, 210, 211; Bayes factors, 220–221; Bayesian information criterion, 211; Bayesian model averaging, 218–219, 223; cross-validation, 207–208; deviance information criterion, 212; effective number of parameters, 212, 215; guidance on use, 223–224; indicator variable selection, 216–218; model probabilities, 218–219, 222; model selection, 206–218; model weights, 221; out-of-sample validation, 206–207; parameter counting for hierarchical models, 212; posterior predictive loss, 215; reversible jump MCMC, 221–223; statistical regularization, 209–210; Watanabe-Akaike information criterion, 213–216; wrong way to calculate, 222
- multinomial distribution, 55
- multiple sources of data, 135–139
- multivariate lognormal distribution, 137
- multivariate normal distribution, 62, 135, 261
- negative binomial distribution, 54, 59, 69, 111, 118, 246
- normal distribution, 56, 115, 132, 198, 249
- notation; bracket, 5, 6; conditional on, 32;
 - conventions, 5; covariate, 5; deterministic model, 6; distribution, 41–42; multiple products, 107; multivariate normal, 261; random effect, 113; random effects, replacing with group level effects, 113
- observation model, 8, 10, 228–230, 232–233
- observation variance, 30, 67, 117, 118, 121, 133, 138, 139, 196, 204, 256; separating from process variance, 139–140
- offset, 53
- OpenBUGS, 4, 312
- out-of-sample validation, 206–207
- over dispersion, 68
- parameter model, *See also* prior distribution, 9
- PDF, *See* probability density function
- peer review of Bayesian models, 1–4, 309–312
- penalty, 210, 212, 216
- Picasso, Pablo, 3

- Poisson distribution, 52–139
- posterior distribution, 83–84, 88, 101, 106, 127, 133, 180, 189, 235, 238; as component of Bayes' theorem, 83, 84; averaging across models, 218, 222; calculating parameters using conjugate relationships, 1, 85–88, 244, 309; influence of prior distribution, 87–201; marginal, 181, 183, 188, 201, 207; of derived quantity, 191–260; of difference between parameters, 192, 251, 260; of population growth rate, 192; of unobserved quantity, 193–198; plotting over prior distribution, 2, 201, 310; predictive, 181, 184, 185, 193–207; when to include predictor variables, 116, 120–121, 128, 242, 251; writing expression for, 111
- posterior predictive check, *See* model checking
- posterior predictive distribution, 180–185, 193–224; of mean, 194; of observation, 194–195
- posterior predictive loss, 215
- power function, 26, 131
- predictions of unobserved quantity from MCMC output, 192–197
- predictive process distribution, 196–197
- predictor variable, *See* covariate
- prior distribution, 88–104; component of Bayes' theorem, 82; conjugate, 85–88, 164, 166; dependent parameters, 104; formulating, 100–102; guidance on use, 102, 104; informative, 96–97, 198, 203, 212; Jefferys prior, 93; joint, 217; non informative, 88, 95; parameters for vague normal distribution, 93; propriety, 91–93; reference, 95; sensitivity analysis, 103; slab and spike, 217; use in maximum likelihood, 76–77; vague, 91, 203; vague priors are provisional, 103; visualizing by plotting, 103; with large variance, 93
- prior predictive distribution, *See* marginal distribution of data
- probability density, 43
- probability density function, 43
- probability distribution, *See* distribution
- probability function, *See* probability mass function
- probability mass function, 40
- probability of data, *See* marginal distribution of data
- probability, definition, 33
- probability, rules of, 31–39; conditional events, 33, 34; disjoint events, 35; factoring joint distributions, 36, 38; independent events, 34
- probability, union of events, 35
- probability: law of total probability, 34
- procedural approach to statistical training, x
- process model, 5–7, 10, 11, 131, 135, 198, 229, 230, 252, 254
- process variance, 30, 115, 116, 118, 120, 131, 137, 139, 196, 204, 248, 249, 251, 255, 258; separating from observation variance, 139–140
- proposal distribution, 149–151, 162; dependent, 149; independent, 149; symmetric, 150; symmetric, 149–151, 155; tuning parameter, 149
- quantile function, 47
- Raftery-Lewis diagnostic, 172
- random draws from posterior distribution, 145
- random effect, *See also* group-level effect, 112, 121, 129–131, 229
- random variable: continuous, 40; definition, 31; discrete, 40; importance to Bayesian inference, 31
- reference prior, 95
- regularization, statistical, 209–210
- regularizer, 209
- regularizer function, 209, 212, 216
- sampling model, 7, 10, 229, 230, 233
- sampling variance, 251

- second central moment, 44
- Serengeti wildebeest model, 10–14, 197–204
- shrinkage, 2, 7, 125, 201, 209, 310
- simple Bayesian model, 84, 87, 106, 109, 258
- simulated data for verifying model coding, 172–174
- slab and spike prior distribution, 217
- software, xiv, 4, 213, 312; BUGS, 175–177; INLA, 2, 310; JAGS, x, 187; OpenBUGS, 4, 312; R, 50; Stan, 2, 310; WinBUGS, x, 4, 187, 312
- spatial models, 289–307; areal data, 294; conditional autoregressive models, 294; continuous spatial models, 291–293; geostatistical model, 293; kriging, 297; simultaneous autoregressive models, 295; temperature example, 295–297
- Stan, 2, 310
- standardizing covariates, 254
- statistical regularization, 209–210
- stochastic search variable selection, 217
- stochasticity, sources, *See* uncertainty
- support, 22, 43, 56, 65; definition, 29, 43
- texts, Bayesian modeling, ix
- thinning MCMC, 168
- threshold model, 28
- trace plot, 169, 170, 200
- trace plots, 170
- uncertainty: observation, 7, 133; individual variation, 30; model selection, 31; partitioning observation and process variance, 139–140; process, 7, 30, 133; sampling, 7, 129
- uniform distribution, 62
- variance: observation, 30, 67, 117, 118, 121, 133, 138, 139, 196, 204, 256; process, 30, 115, 116, 118, 120, 131, 137, 139, 196, 204, 248, 249, 251, 255, 257; sampling, 233, 251
- WAIC, *See* Watanabe-Akaike information criterion
- Watanabe-Akaike information criterion, 213–215
- WinBUGS, x, 4, 187, 312
- zero inflation, 68, 239, 252